



Universiteit  
Leiden  
The Netherlands

## Orchestration of Distributed LOFAR Workflows

Mechev, A.P.

### Citation

Mechev, A. P. (2019, December 9). *Orchestration of Distributed LOFAR Workflows*. Retrieved from <https://hdl.handle.net/1887/81487>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/81487>

**Note:** To cite this publication please use the final published version (if applicable).

Cover Page



Universiteit Leiden



The following handle holds various files of this Leiden University dissertation:  
<http://hdl.handle.net/1887/81487>

**Author:** Mechev, A.P.

**Title:** Orchestration of Distributed LOFAR Workflows

**Issue Date:** 2019-12-09



# Orchestration of Distributed LOFAR Workflows

## PROEFSCHRIFT

ter verkrijging van  
de graad van Doctor aan de Universiteit Leiden  
op gezag van de Rector Magnificus Prof. mr. C.J.J.M. Stolker,  
volgens besluit van het College voor Promoties  
te verdedigen op maandag 9 December 2019 klokke 15:00 uur

door  
Alexandar Plamenov Mechev  
geboren te Sofia, Bulgarije  
in 1989

Promotiecommissie:

Promotor: Dhr. Prof.dr. H.J.A. Röttgering  
Promotor: Dhr. Prof.dr. A. Plaat  
Co-Promotor: Dr. J.B.R. Oonk

Overige leden: Dhr. Prof.dr. S.F. Portegies Zwart  
Dhr. Prof.dr. H.A.G. Wijshoff  
Prof.dr. Rob van Nieuwpoort  
Prof.dr. Martin Hardcastle  
Dr. A.L. Varbanescu  
Dhr. Dr.ing. H.T. Interna  
Dhr. Dr. T.W. Shimwell

The cover of this thesis depicts a parent duck orchestrating a neat line of ducklings, inspired by a common sight during the author's commute by bike. The theme of this work, orchestration of complex workflows, can be expressed by the idiom "having your ducks in a row", and the cover page illustrates this.

# Contents

---

|          |                                                    |           |
|----------|----------------------------------------------------|-----------|
| <b>1</b> | <b>Introduction</b>                                | <b>1</b>  |
| 1.1      | Introduction . . . . .                             | 1         |
| 1.2      | LOFAR . . . . .                                    | 6         |
| 1.3      | Problem Statement and Research Questions . . . . . | 12        |
| 1.4      | Contributions . . . . .                            | 14        |
| <b>2</b> | <b>LOFAR-DSP platform</b>                          | <b>17</b> |
| 2.1      | Introduction . . . . .                             | 17        |
| 2.2      | LOFAR observations . . . . .                       | 19        |
| 2.3      | LOFAR Processing . . . . .                         | 21        |
| 2.4      | LOFAR-DSP . . . . .                                | 23        |
| 2.5      | Executing LOFAR-DSP . . . . .                      | 29        |
| 2.6      | Deploying LOFAR-DSP . . . . .                      | 32        |
| 2.7      | Discussion . . . . .                               | 35        |
| 2.8      | Conclusions . . . . .                              | 38        |
| <b>3</b> | <b>LOFAR Scalability Framework</b>                 | <b>43</b> |
| 3.1      | Introduction . . . . .                             | 44        |
| 3.2      | Related Work . . . . .                             | 45        |
| 3.3      | LOFAR Data Processing . . . . .                    | 46        |
| 3.4      | Framework Design . . . . .                         | 49        |
| 3.5      | Conclusion and Future Work . . . . .               | 54        |
| 3.A      | Execution of LOFAR Reduction Tools . . . . .       | 57        |

|          |                                                                        |            |
|----------|------------------------------------------------------------------------|------------|
| <b>4</b> | <b>Pipeline Collector</b>                                              | <b>59</b>  |
| 4.1      | Introduction . . . . .                                                 | 60         |
| 4.2      | Measuring LOFAR Pipeline performance with pipeline_collector . . . . . | 63         |
| 4.3      | LOFAR Prefactor Test Case . . . . .                                    | 66         |
| 4.4      | CPU Utilization Tests with PAPI . . . . .                              | 73         |
| 4.5      | Discussions and Recommendations . . . . .                              | 77         |
| 4.6      | Conclusions . . . . .                                                  | 79         |
| 4.A      | Performance Collection Implementation Details . . . . .                | 80         |
| <b>5</b> | <b>Fast and Reproducible LOFAR Workflows with AGLOW</b>                | <b>83</b>  |
| 5.1      | Introduction . . . . .                                                 | 84         |
| 5.2      | Background . . . . .                                                   | 85         |
| 5.3      | Related Work . . . . .                                                 | 86         |
| 5.4      | AGLOW . . . . .                                                        | 87         |
| 5.5      | Results and Discussions . . . . .                                      | 96         |
| 5.6      | Conclusions . . . . .                                                  | 97         |
| <b>6</b> | <b>Scalability Model for the LOFAR Direction Independent Pipeline</b>  | <b>101</b> |
| 6.1      | Introduction . . . . .                                                 | 102        |
| 6.2      | Related Work . . . . .                                                 | 104        |
| 6.3      | Processing Setup . . . . .                                             | 104        |
| 6.4      | Results . . . . .                                                      | 109        |
| 6.5      | Discussions and Conclusions . . . . .                                  | 122        |
| 6.6      | Applications and Conclusions . . . . .                                 | 127        |
| 6.A      | Calibration Solutions for the sky model tests . . . . .                | 128        |
| 6.B      | Parametric model parameters and fit accuracy . . . . .                 | 129        |

|                                                                                         |            |
|-----------------------------------------------------------------------------------------|------------|
| <b>7 Automated testing and quality control of LOFAR scientific pipelines with AGLOW</b> | <b>133</b> |
| 7.1 Introduction . . . . .                                                              | 133        |
| 7.2 Background . . . . .                                                                | 136        |
| 7.3 Related Work . . . . .                                                              | 137        |
| 7.4 Automated testing with AGLOW . . . . .                                              | 138        |
| 7.5 Results . . . . .                                                                   | 141        |
| 7.6 Discussion and Conclusions . . . . .                                                | 145        |
| <b>8 Conclusion</b>                                                                     | <b>149</b> |
| 8.1 Summary of Thesis Contributions . . . . .                                           | 149        |
| 8.2 Answers to Research Questions . . . . .                                             | 150        |
| 8.3 Limitations . . . . .                                                               | 151        |
| 8.4 Future Work . . . . .                                                               | 152        |
| <b>Acknowledgements</b>                                                                 | <b>155</b> |
| <b>Samenvatting</b>                                                                     | <b>159</b> |
| <b>English Summary</b>                                                                  | <b>165</b> |
| <b>List of Publications</b>                                                             | <b>169</b> |
| 8.5 Journal Articles . . . . .                                                          | 169        |
| 8.6 Conference Proceedings . . . . .                                                    | 170        |
| <b>Bibliography</b>                                                                     | <b>172</b> |



## List of Tables

---

|     |                                                                          |     |
|-----|--------------------------------------------------------------------------|-----|
| 2.1 | Comparison of LOFAR software distribution methods. . . . .               | 27  |
| 4.1 | A table of all the results presented in Section 4.3. . . . .             | 62  |
| 4.2 | Hardware specifications of the four test machines. . . . .               | 66  |
| 6.1 | Averaging parameters and final data sizes for a sample LOFAR Observation | 106 |
| 6.2 | List of test sky models . . . . .                                        | 107 |
| 6.3 | Image statistics for four different sky models . . . . .                 | 115 |
| 6.4 | Queueing statistics per requested number of CPUs . . . . .               | 118 |
| 6.5 | Fit parameters for the models in Equation 6.1. . . . .                   | 131 |
| 6.6 | Goodness of fit parameters for the model in Equation 6.4. . . . .        | 132 |



## List of Figures

---

|     |                                                                                                       |    |
|-----|-------------------------------------------------------------------------------------------------------|----|
| 1.1 | Two supercomputers sixty years apart. . . . .                                                         | 2  |
| 1.2 | Graphical representation of aperture synthesis . . . . .                                              | 5  |
| 1.3 | Image of the raw data . . . . .                                                                       | 10 |
| 1.4 | Image of preprocessed data . . . . .                                                                  | 10 |
| 1.5 | Image of DI calibrated data . . . . .                                                                 | 11 |
| 1.6 | Fully calibrated image . . . . .                                                                      | 11 |
| 2.1 | Structure of the LOFAR-DSP platform. . . . .                                                          | 40 |
| 2.2 | Implementation of Pilot Job Launching . . . . .                                                       | 41 |
| 3.1 | Prefactor and DDFacet Data Flow . . . . .                                                             | 46 |
| 3.2 | Parallelization of prefactor processing . . . . .                                                     | 49 |
| 3.3 | Overview of the design of the LRT framework. . . . .                                                  | 50 |
| 3.4 | Starting processing on worker machines . . . . .                                                      | 52 |
| 3.5 | Schematic of data movement. . . . .                                                                   | 53 |
| 4.1 | The four processing stages that make up the prefactor pipeline. . . . .                               | 64 |
| 4.2 | Portion of processing time taken by each step for the four prefactor stages. . . . .                  | 65 |
| 4.3 | Job completion times for two processing steps tested on four hardware setups . . . . .                | 68 |
| 4.4 | Speed comparison between natively and remotely compiled software for the<br>'calib_cal' step. . . . . | 69 |

|      |                                                                                                                                                                            |     |
|------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 4.5  | Speed comparison between natively and remotely compiled software for the 'gsmcal_solve' step. . . . .                                                                      | 69  |
| 4.6  | A model of the memory hierarchy, as described in [104]. . . . .                                                                                                            | 70  |
| 4.7  | Effect of CPU speeds on the bottle neck steps for the four test machines. . .                                                                                              | 71  |
| 4.8  | Effect of cache size on the bottle neck steps for the four test machines. . . .                                                                                            | 72  |
| 4.9  | Effect of RAM throughput on the bottle neck steps for the four test machines.                                                                                              | 73  |
| 4.10 | Time series of the Virtual Memory Resident Set Size . . . . .                                                                                                              | 74  |
| 4.11 | Performance of the two bottleneck steps and Disk bandwidth in MB/s. . . .                                                                                                  | 74  |
| 4.12 | Cache miss rates for the bottleneck steps, executed on the SURFsara gina cluster. . . . .                                                                                  | 75  |
| 4.13 | Resource stall cycles and Full Instruction Issue cycles. . . . .                                                                                                           | 76  |
| 4.14 | Communication between worker nodes and the TSDB server, including the pipeline_collector modules (in red). . . . .                                                         | 81  |
| 5.1  | Design of the AGLOW software, including its constituent packages. . . . .                                                                                                  | 89  |
| 5.2  | Graphical representation of the Airflow Operators built for AGLOW . . . .                                                                                                  | 92  |
| 5.3  | Rendering of the DAG encoding a DPPP parset as shown by the Airflow User Interface. . . . .                                                                                | 93  |
| 5.4  | Render of the DAG encoding the full prefactor pipeline. . . . .                                                                                                            | 98  |
| 6.1  | The major steps of the <code>prefactor</code> DI pipeline. . . . .                                                                                                         | 106 |
| 6.2  | The size of the sky model (measured as the number of sources) increases exponentially as we decrease the flux cutoff of the model (i.e. increase the sensitivity). . . . . | 108 |
| 6.3  | Plots of the run time as a function of input data size . . . . .                                                                                                           | 111 |
| 6.4  | Tests of the <code>gsmcal_apply</code> step for data from 1GB to 64GB . . . . .                                                                                            | 112 |
| 6.5  | Processing time of <code>gsmcal_solve</code> vs number of sources in <code>skymodel</code> . . .                                                                           | 113 |
| 6.6  | Run time vs cutoff sensitivity . . . . .                                                                                                                                   | 114 |
| 6.7  | Comparison of data calibrated with four sky models . . . . .                                                                                                               | 114 |
| 6.8  | The processing time of the <code>gsmcal_solve</code> step vs number of CPUs requested                                                                                      | 116 |

|      |                                                                                                                  |     |
|------|------------------------------------------------------------------------------------------------------------------|-----|
| 6.9  | No speed-up for <code>gsmcal_apply</code> seen . . . . .                                                         | 116 |
| 6.10 | Queueing tests on the <code>GINA</code> cluster . . . . .                                                        | 117 |
| 6.11 | Model of queueing times . . . . .                                                                                | 119 |
| 6.12 | Histogram of download times . . . . .                                                                            | 119 |
| 6.13 | Histogram of extraction times . . . . .                                                                          | 120 |
| 6.14 | Exponential model of data transfer/extraction on the <code>GINA</code> cluster . . . . .                         | 120 |
| 6.15 | Tet of download/extract times for ten 1GB data sets . . . . .                                                    | 121 |
| 6.16 | Tet of download/extract times for a 64GB data set . . . . .                                                      | 122 |
| 6.17 | Processing time for the <code>gsmcal_solve</code> step in a production environment . . . . .                     | 123 |
| 6.18 | Comparison of scalability model with production runs . . . . .                                                   | 124 |
| 6.19 | Calibration Solutions for sky models with low flux cutoff . . . . .                                              | 129 |
| 6.20 | Calibration solution differences between two skymodels . . . . .                                                 | 130 |
| 6.21 | Calibration solutions for sky models with high flux cutoff . . . . .                                             | 130 |
| 7.1  | Diagram of the <code>prefactor</code> Continuous Integration workflow . . . . .                                  | 140 |
| 7.2  | A test image created on 2019-04-23 by our automated CI workflow showing<br>a diffuse radio source . . . . .      | 142 |
| 7.3  | Diagram of integration of three scientific pipelines with test data and pro-<br>cessing infrastructure . . . . . | 142 |
| 7.4  | Images created by the CI runs from 2019-03-29 and 2019-04-23 showing<br>a bright source . . . . .                | 143 |
| 7.5  | Time series of measurements done with images automatically produced by<br>our workflow . . . . .                 | 143 |



## 1.1 Introduction

For almost a century, scientists have used computational machines as a tool to conduct scientific research. In part driven by national security concerns during the First and Second World Wars, as well as the Cold War, increasingly complex computers have been designed. Early computers were entirely designed for application-specific tasks, with a considerable drive behind them being hydrodynamics simulations for the first Hydrogen bomb. In 1945, the popular Von-Neumann[1] architecture was developed to make Monte Carlo simulations easier to develop and to facilitate general-purpose computing. This architecture was a significant improvement over previous computers where changing the program required physically flipping switches and changing cables on the computer itself. One of the first computers built according to the von-Neumann architecture was the MANIAC[2] computer commissioned by Los Alamos National Laboratory, seen on the left in Figure 1.1.

With the advancement of an architecture that treats code and data identically, it was possible to create more complex programs including compilers: programs that could create machine code from human-readable code. As the '50s and '60s passed, general-purpose computers were increasingly used in science. From weather dynamics[3] to fluid dynamics[4], from chaos theory to game theory, these computers were being adopted by a wide range of scientific fields. Astronomy was, likewise, also a driving force for computational innovation. In 1953, for example, the first high-level programming language for IBM computers was developed by John Backus, a programmer frustrated with the difficulty of accurately calculating the moon's position using only machine code. John Backus' 'Speedcode' was a direct predecessor of Fortran[5], a language developed at IBM in the '50s and still used by the scientific community today. Another important discovery on our road was

the Fast Fourier Transform (FFT), discovered by two researchers from Princeton and IBM[6]. The FFT has been described as ‘the most important numerical algorithm of our lifetime’ and the author’s personal favourite ‘an algorithm the whole family can use’[7]. As we will soon see, Radio Astronomers quickly became part of this family.

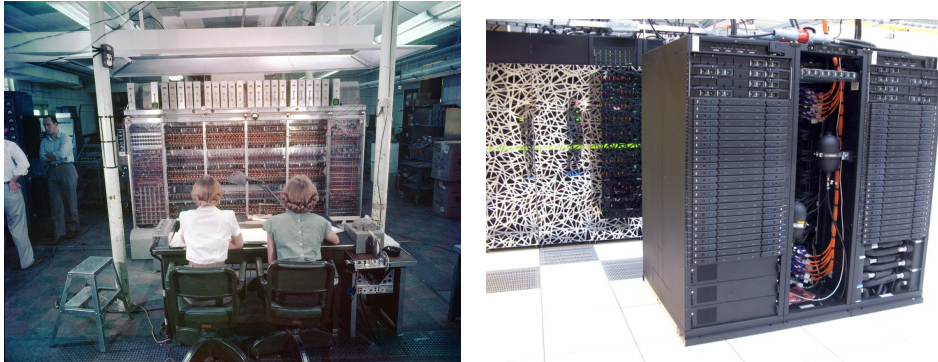


Figure 1.1: Two supercomputers sixty years apart. On the left is the MANIAC computer from 1952 at Los Alamos, while on the right is the Cartesius cluster at SURFsara, Amsterdam in 2018.

As computers became more widely available, they became increasingly adopted by universities and research institutes. In the '70s, computers began to talk to each other over a network connection. This capability not only made scientific collaboration easier but also made it possible to distribute computation across multiple sites. Moreover, the development of the integrated circuit and subsequent drop in price/performance of computers made it financially feasible for scientific institutes to purchase multiple computers dedicated to scientific research. As hardware, networks, and software matured, clusters of computers became more widely used[8]. In part because of their cost-effectiveness, and potential for parallelization, computer clusters became more widely used as the '80s wound down. By then, general-purpose computing was widely adopted by the astronomical community. In the '80s several astronomical software suites have been developed, with software such as AIPS[9] and IRAF[10] and standards such as FITS[11] used to this day.

The '90s continued the distributed computing trend with the appearance of commodity compute clusters, with virtual ‘supercomputers’ being created from Commercial-Off-The-Shelf (COTS) hardware, and networking. These clusters became quickly adopted by scientists to perform simulations and data processing. Around the same time, the idea of ‘grids’ was created[12]. The concept was that of a country, continent, or even worldwide network of hardware that can trans-

parently handle distributed tasks, and provide researchers with a vast pool of resources. CERN pioneered Grid processing to meet the computational and storage requirements of their High Energy Physics modelling and data reduction. While this infrastructure was built initially for HEP experiments, it is also useful for other scientific projects, particularly low-frequency Radio Astronomy.

### 1.1.1 Astronomy and Computing

Since the early days of computing, the field of astronomy has embraced digitization of data acquisition and processing. Being able to store astronomical data digitally makes it possible to transfer, copy, backup, and process them efficiently. While optical astronomy entered the digital age in the '80s with the rapid development of CCDs, thanks to the extensive availability of Analog-Digital Converters (ADCs), radio astronomy has been digital since the 1970s. By the end of the '70s, the Very Large Array (VLA) in New Mexico and the Westerbork Synthesis Radio Telescope (WSRT) had consistently been using processing pipelines, running on IBM mainframes, Digital Equipment Corporation's line of PDP, and later VAX, minicomputers. Notably, their imaging algorithms were taking advantage of the FFT developed a decade earlier[13].

With the complete digitization of astronomical observations, over the past decade, all of astronomy has entered the big data regime. As of 2019, there are multiple planned and ongoing large-scale sky surveys across the electromagnetic spectrum, each expecting to produce multiple tens of petabytes. This breadth of data is poised to expand the frontiers of astronomy and astrophysics and allow us to study and understand various phenomena in more detail.

The longest wavelength of the spectrum accessible to Earth observatories lies in the Megahertz range, starting at 10 MHz, up to 300 MHz. This regime corresponds to wavelengths of 30 meters up to 1 meter. In astronomy, this range is termed the low-frequency or meter-wave regime. These wavelengths help uncover physical phenomena invisible to telescopes in the X-ray, Visible, Infrared, or Microwave. In particular, long-wavelengths can be used to study supermassive black holes, galaxy formation and evolution, magneto-hydrodynamics, solar physics, radio spectroscopy, and many more science cases. Additionally, data in this domain can complement other telescopes in multi-wavelength studies.

Photons provide us with rather few independent properties. We can use them to record the wavelength, the direction, the intensity, and the polarization of

a distant source. We can, moreover, record the change of those properties with time. Astronomers need to measure properties accurately and use the data to model distant sources better and validate or reject astronomical theories. The accuracy of these models, or the rejection power of our observations, depends critically on how accurately we can measure the four properties listed above.

Astronomical observations in the long-wavelength regime have always been at the mercy of the diffraction limit, an effect that relates the wavelength of light, the diameter of the aperture, and angular resolution obtained with that aperture. The angular resolution of a telescope determines how accurately the direction of an incoming photon can be determined. Unfortunately, the diffraction limit dictates that the angular resolution of a telescope with a fixed aperture decreases inversely proportional to the wavelength observed. For example, if one takes a telescope at 100MHz and one at 10GHz, the 100MHz telescope would need to have 100 times the radius of its higher frequency counterpart in order to reach the same angular resolution. In other words, for a low-frequency telescope (at 100 MHz) to match the 100-m Effelsberg telescope (at 10GHz), it would need a dish with a diameter of 10 kilometres. Constructing, and operating a telescope of that size is currently outside our engineering capabilities, and thus low-frequency astronomers have developed a method to synthesize a telescope aperture of arbitrary size, termed ‘Aperture Synthesis.’

### 1.1.2 Aperture Synthesis

Aperture synthesis is the practice of combining the signal of multiple antennas to produce data with the angular resolution of a much larger antenna, as seen in Figure 1.2. More specifically, the maximum angular resolution achievable is proportional to the distance between the furthest two antennas. This technique is used in a wide wavelength range, from the near- and mid-infrared (e.g., VLTI), sub-millimeter (e.g., ALMA) and radio wavelengths (e.g., VLA, GMRT, LWA).

Aperture synthesis can be used to increase the angular resolution of individual radio telescopes; however the resulting data requires significant post-processing. A single telescope operates in the ‘image’ domain, meaning the data collected at its focus is directly related to an area on the sky. Conversely, an array of telescopes records correlated signals between pairs of antennas. We need to transform this data to obtain a map of the radio sky. The equation relating the data recorded by these arrays and the ‘true’ sky distribution is the Radio Interferometry Measurement Equation, RIME. This equation describes propagation effects from the source  $B$  to

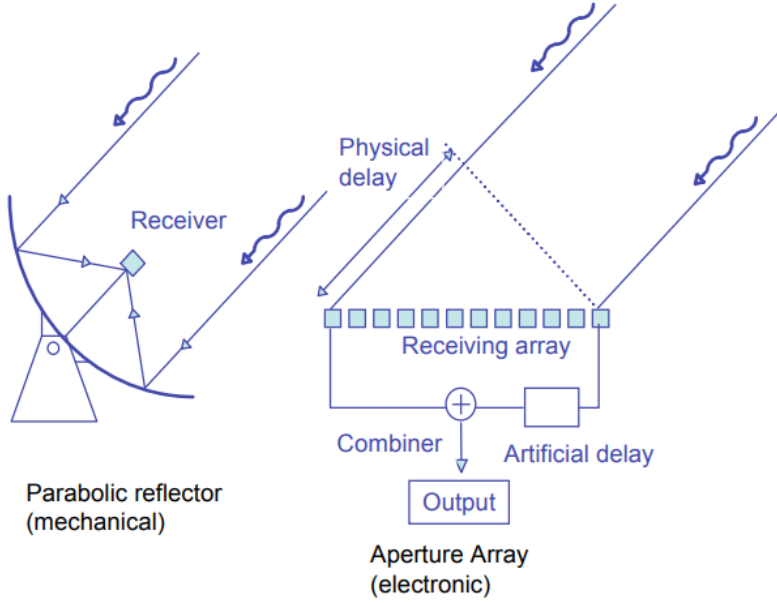


Figure 1.2: We can simulate a single dish with an array of antennas. These antennae are pointed in different directions by introducing a corresponding hardware delay in each antenna feed. While this process can enable us to synthesize an arbitrarily large telescope, it produces artefacts in the final image that need extensive processing to remove.

two antennas,  $p$  and  $q$ , with  $n$  effects towards antenna  $p$  and  $m$  effects towards antenna  $q$ . Each effect is described by a  $2 \times 2$  ‘Jones’ matrix describing the transformation of the original signal. This formulation is shown in Equation 1.1. When expressed in terms of the directions  $(l, m)$ , and a continuous sky model,  $B$ , the measurement equation becomes Equation 1.2[14] .

Equation 1.2 neatly separates the direction-independent terms ( $\mathbf{G}_p$  and  $\mathbf{G}_q^H$ ) seen by antenna  $p$  and  $q$ , and the direction-dependent effects that correspond to directions  $l$  and  $m$  inside the integral. A comparison between equations 1.2 and 1.3 shows the similarity between the RIME and the Fourier Transform. Specifically,  $f(x)$  represents the sky brightness  $B$ ,  $\xi$  represents our directions  $l$  and  $m$ , and the transformed function  $f(\xi)$  are the visibilities ( $V_{pq}$ ) measured by the telescope (Equation 1.3). This formalism also separates the direction-dependent and independent effects, and further shows how we can use Fourier transforms to obtain a model of the radio sources. As Fourier Transforms are computationally expensive, an efficient

solution is to use the Fast Fourier Transform algorithm mentioned above.

An aperture synthesis telescope consisting of  $N$  antennas will also have  $\frac{N^2-N}{2}$  baselines, each baseline being defined by a unique pair of antennas. The length and orientation of each of these baselines corresponds to a single location in Fourier space. Earth's rotation helps sampling the Fourier space by changing the (projected) baseline length and orientation with each time sample[15]. Because the telescope doesn't completely fill the Fourier space, its point spread function (PSF) will have large, extended side-lobes. The final image is convolved with this PSF, creating artifacts such as those seen in Figures 1.3 and 1.5. Removing these artifacts from raw data requires estimating gain parameters by LM-minimization and fitting[16], followed by multiple cycles of convolutions, subtractions and deconvolutions[17–19].

$$\mathbf{V}_{pq} = \mathbf{J}_{pn}(\dots(\mathbf{J}_{p2}(\mathbf{J}_{p1}\mathbf{B}\mathbf{J}_{q1}^H)\mathbf{J}_{q2}^H)\dots)\mathbf{J}_{qm}^H \quad (1.1)$$

$$\mathbf{V}_{pq} = \mathbf{G}_p \left( \iint_{lm} \frac{1}{n} \mathbf{E}_p \mathbf{B} \mathbf{E}_q^H e^{-2\pi i(u_{pq}l + v_{pq}m + w_{pq}(n-1))} dldm \right) \mathbf{G}_q^H \quad (1.2)$$

$$\hat{f}(\xi) = \int f(x) e^{-2\pi i x \xi} dx \quad (1.3)$$

In this work, we will discuss the technical challenges of creating radio images of astronomical sources and our solutions to these challenges. In the following section, we aim to introduce LOFAR, the European Low-Frequency Array, the data sizes and processing challenges that come with LOFAR data as well as our solutions to these challenges. We will conclude with the scientific results this work has led to, as well as suggestions for future large-scale astronomical projects.

## 1.2 LOFAR

LOFAR is a large low-frequency radio telescope centered near Dwingeloo, Drenthe, in the Netherlands[20]. It is designed to operate between 10 and 240 MHz, but it cannot observe in the FM radio bands from 80-120 MHz. Thus, LOFAR is split into two arrays: the Low Band Array (LBA) operating from 10 to 80 MHz and the High Band Array (HBA) with antennas sensitive to frequencies from 120 to 240

MHz. In the Netherlands alone LOFAR has more than 5000 antennas: 1824 High-Band antennas and 3648 Low-Band antennas. These are grouped in core (near Dwingeloo) and remote stations[21]. Additionally, LOFAR has 13 international stations across Europe, spanning from Ireland to Latvia, Sweden to France[22]. These international stations make it possible to create images of radio sources with a similar angular resolution to leading higher frequency telescopes. LOFAR was also designed to support a variety of science cases, such as studying the Epoch of Reionization[23], performing large scale extragalactic surveys[24–26], studying cosmic magnetism, radio spectroscopy[27–30] and transient detection[31, 32].

LOFAR stores its broadband observations at one of several Long-Term Archive locations. These locations store the data on tape, due to its large size and infrequent access. Typical broadband observations are up to 16TB in size, which can drop down to 10TB with compression. While individual researchers use this data to study their object of interest, the majority of the broadband data will be imaged to produce the LOFAR Two-Meter Sky Survey (LoTSS).

### 1.2.1 LoTSS

The LOFAR Two-Meter Sky Survey, LoTSS[25], is an ambitious project to map the Northern Radio sky at low frequencies, namely 120-168 MHz. Expected to produce more than 3000 8-hour observations, LoTSS will create radio maps with a median sensitivity of  $70 \mu\text{Jy}/\text{beam}$ . This survey will help study supermassive black holes and their impact on galaxy formation in the early Universe. Additionally, addressing questions related to the formation and evolution of galactic clusters and the interaction of galaxies within these clusters will be made possible with this low-frequency data. Furthermore, the survey will enable us to study star formation in nearby and distant galaxies and galactic sources such as supernova remnants. Finally, LoTSS will help study and discover patterns in the large-scale structure of the Universe.

#### Processing Requirements

With its 3000+ observations, the LoTSS project requires a large amount of processing, bandwidth, and storage infrastructure in order to complete its scientific goals within the survey timespan. The total size of raw data is more than 30 petabytes, while the total size of the finished products will be on the order of 10s of terabytes.

Furthermore, moving all the raw data to a processing facility is limited by the bandwidth of the connection between the archive site and the processing facility. Finally, producing a high fidelity image from each data set requires roughly 3500 core-hours. In total, this means that the LoTSS project will take more than 10 million core-hours to produce scientific results, assuming no re-processing of data.

In addition to the raw hardware requirements, an ambitious project as such needs to be able to track the status and location of data products, automate processing and make results readily available. As multiple locations store LOFAR data, it is also essential that the framework tasked with processing LoTSS data is portable and can run independently of the infrastructure details.

### 1.2.2 SURFsara

One of the archive locations storing LOFAR data is SURFsara at the Amsterdam Science Park. Aside from an extensive storage archive, SURFsara also supports several clusters, including the Gina cluster, part of the Dutch Grid infrastructure. Grid computing is a non-interactive application-oriented computational paradigm for distributed computing where a 'grid' consists of a large pool of nodes where users can submit batch jobs. A grid can consist of one cluster or groups of clusters at one or multiple geographical locations, connected with high-speed links and a standard job management interface. Using this interface, users can scale out their projects, given that their processing is massively parallel. Computational resources on such a platform are granted based on the quality of a scientific proposal and are used freely across the Grid, while jobs are scheduled based on the job requirements and the current resource availability of the grid nodes. This processing paradigm is perfect for extensive grid-search simulations, but also the first steps of LOFAR processing. Furthermore, the high-speed connection to the LOFAR archive and available storage makes SURFsara a logical location to orchestrate large scale LOFAR projects and distribute processed data.

### 1.2.3 LoTSS Processing

The observations produced by LoTSS will total more than 30 petabytes and require extensive processing before completing the survey. Each observation is stored as a set of 244 individual files spanning frequency space from 120MHz to 168MHz. Each of these files, named a Subband, is a CASA Measurement set[33] and is identified by its three-digit Subband number, starting from 000. Each Subband, thus,

contains a sub-sample of the data in frequency space, stored at a resolution of 1 second and 12.2 kHz per sample. While this high-resolution data is useful for some science cases, our processing algorithms scale with data size, and thus it is necessary to average our data in order to complete the LoTSS processing within the project's time-frame.

In order to create an image from an archived data set, the data needs to be staged, retrieved, and processed. Staging the data refers to sending a request to the archive site to move the data from tape to disk. Once all the data is on disk ('staged'), it is ready to be transferred from the storage to the processing cluster. On this cluster, a science-ready image is produced by processing the raw data through two pipelines. The first pipeline, Direction-Independent Calibration pipeline removes artifacts created by 'direction-independent' effects, i.e. effects that are constant across the field of interest. This pipeline is followed by the Direction-Dependent Calibration pipeline, which removes effects that change within the field of view.

The Direction-Independent Calibration pipeline (DI pipeline) consists of two main stages. The first stage is calibration on the calibrator, which uses a short observation of a bright calibration source to determine systematic effects that are independent of the direction of the pointing. These effects include phase errors due to station clock offsets and Direction Independent ionospheric corrections. The solutions obtained from this step can be applied to the scientific target, improving the data quality. The second step of the DI pipeline is the calibration of the target field against a sky model produced by a previous survey. This calibration determines the gain parameters of all antennas; however it does not correct for effects that vary across the field of view.

In order to create a high fidelity radio image, we need to correct for effects that not only change in time but also across the field of view. These effects, such as the ionosphere or the beam response can be modelled and removed, and their removal is the responsibility of the Direction-Dependent pipeline (DD pipeline). Upon successful completion, the DD pipeline produces a radio image to be used for further scientific studies.

There are a few software packages used to process LOFAR data, typically used in a series of steps creating a processing pipeline. The LOFAR processing pipeline steps use the software, each step encoding the processing parameters in a parameter-set (parset). A pipeline is defined by the list of steps, along with the parameters of each step concatenated together into a parset file. The LoTSS DI pipeline, `prefactor` contains a set of scripts used to remove direction-independent

effects from LOFAR data. Many of the prefactor steps can be executed on the data in parallel: each Subband can be processed independently. Because of the large amount of data, the best architecture for these steps is a cluster of isolated machines with dedicated disks and a high-speed connection to the data. Later we will show the benefit of automating these steps on the Dutch Grid infrastructure.

### 1.2.4 The life of a data set

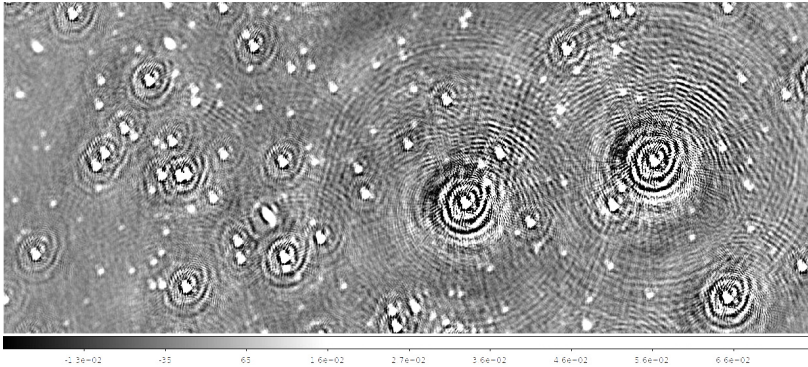


Figure 1.3: Raw data for LOFAR observation L229587. We only image half of the bandwidth, from Subband 061 to Subband 183. This data was retrieved directly from the LOFAR archive and thus has minimal corrections applied to it. The bright rings around most sources are an indication that the data has not been calibrated to remove Direction Independent (DI) or Direction Dependent (DD) effects.

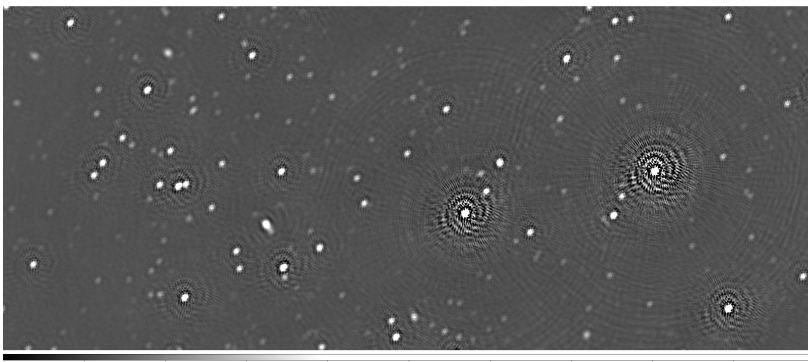


Figure 1.4: Preprocessed data for the same observation as figure 1.3. This image was produced by applying the calibrator solutions to the raw data. These solutions are generated from a short observation of a bright standard calibration source. Note the corrected scale bar, and the reduced prominence of artifacts around most sources.

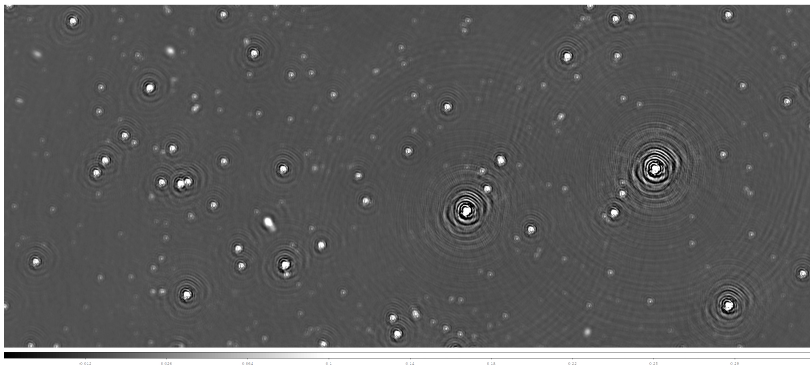


Figure 1.5: Data from L229587 after calibration against a global sky model. This model, obtained from a survey by a previous telescope, is used to determine the antenna phases that result. This calibration removes the direction-independent effects.

In this section, we will show the progress of one observation<sup>1</sup> from raw data to a final scientific image. We run this observation through the `prefactor` pipeline<sup>2</sup> to perform the Direction Independent correction, followed by the `ddf-pipeline`<sup>3</sup> for Direction Dependent corrections. For this work, we only use half of the bandwidth (from 132.2MHz to 156.1MHz) to speed up the DD calibration.



Figure 1.6: Full direction-dependent calibrated image of the L229587 data, done at a high resolution. This image shows a drastic reduction in imaging artifacts as well as a low image noise level. Some of the leftover artifacts are due to skipping the flux bootstrapping step, and using half the bandwidth

Figure 1.3 shows an image of the data downloaded from the LOFAR Long Term Archive. There have been minimal corrections done to this data, only re-

<sup>1</sup>The target is P18Hetdex03, observed on 2014-05-28 with phase centre 11h55m41.282, +049d44m52.908

<sup>2</sup>`prefactor` v3.0 beta 1.

<sup>3</sup><https://github.com/mhardcastle/ddf-pipeline>

moving radio frequency interference before archiving the observation. The lack of phase and amplitude corrections result in a large amplitude offset (see scale bar below figure) and distinct artifacts around bright sources. In order to decrease these artifacts and calibrate the brightness of the sources, we use calibration data from a known bright radio source and apply these solutions to our data. The resulting data produces Figure 1.4. We remove Radio Frequency Interference (caused by man-made sources) from our data and correct for bright off-axis radio sources. Finally, we use a model of the radio sky obtained by a previous survey to calibrate all our direction-independent gain parameters. The resulting data produces Figure 1.5.

After the Direction-Independent calibration, we perform a correction for Direction-Dependent effects. This correction is done by the `ddf`-pipeline scripts using the `DDFacet` and `killMS` software packages. To produce the images, we use the ‘tier1-jul2018’ parameters. To speed up processing, we turn off the bootstrapping step, which loads flux estimates from previous surveys[34]. Once all the DD calibration completes, it produces a high resolution, high fidelity image of the target observation, shown in Figure 1.6.

### 1.3 Problem Statement and Research Questions

Radio Astronomy data sets are too large to process in bulk on individual workstations and often strain the resources of small clusters at universities and other institutions. This limitation in resources requires high throughput processing capability, and automation in order to serve processed data in bulk to astronomers. The LOFAR radio telescope acquires data at a rate of roughly a terabyte per hour. This data is stored in a Long Term Archive as it can serve multiple science cases. Our goal is to create the tools for scientists to be able to process this data with their scripts and software efficiently. As such, these tools need to be fast, easy to use, general, and scalable.

***Problem Statement:*** *How can we efficiently process broadband LOFAR data in a generic way?*

#### 1.3.1 Research Question 1

To efficiently process LOFAR data, we need to take advantage of the data level parallelism of the early processing steps, each of which can be parallelized by a factor of 244. Because of the large sizes of LOFAR observations, transfer time can be comparable to the processing time. To minimize transfer time, we need to study how to

deploy processing pipelines at the LTA storage sites. In this research question, we ask how to best build a framework for a massively distributed shared platform for LOFAR, and how to deploy LOFAR processing in parallel when possible.

**Research Question 1:** *How can we use a distributed shared infrastructure for efficient LOFAR data processing?*

### 1.3.2 Research Question 2

Once we have determined the utility of distributed processing for the LOFAR case, we ask how to automate complex LOFAR workflows. The LOFAR radio telescope serves multiple science cases, each of which is served by a multi-step pipeline with a broad set of parameters. Running an entire pipeline on a single computational node is inefficient; thus, a workflow orchestration software is needed to parallelize the appropriate steps. In this research question, we ask how to build software to efficiently integrate scientific pipelines with a massively parallel distributed processing platform.

**Research Question 2:** *How can we build software to effortlessly accelerate complex pipelines for Radio Astronomy?*

### 1.3.3 Research Question 3

Once complex pipelines can be executed on a distributed environment, researchers may ask whether the software concurrently running on hundreds of systems is running optimally. Manually monitoring automated runs is not possible, hence software is needed to collect this performance data per pipeline step. Furthermore, some of the processing parameters for LOFAR pipelines result in large data sets. In order to serve LOFAR processing to the scientific community, we need to understand how our resource usage scales with each of the processing parameters. We ask whether it is possible to integrate monitoring tools to our processing framework in a way that we can transparently collect performance data along with scientific processing.

**Research Question 3:** *Can we automatically collect performance information during massively distributed processing and predict run times for future data sets?*

## 1.4 Contributions

### 1.4.1 Software Contributions

For this work, we built several software packages to define, launch, and orchestrate jobs on a high throughput cluster. The software packages built are GRID\_LRT<sup>4</sup>, GRID\_PiCaS\_Launcher<sup>5</sup>, and AGLOW<sup>6</sup>. They are available on GitHub, and their documentation is hosted on ReadTheDocs.

### 1.4.2 Statement of Originality

I hereby certify that the content of this thesis is my own original work, consisting of six manuscripts submitted to peer-reviewed journals and conferences. This thesis and the works therein have not been submitted to any other degree program. Finally, I certify that the intellectual content in this work and all software referenced therein are of my own work unless explicitly cited otherwise and that all assistance in compiling this work has been adequately acknowledged.

### 1.4.3 Results

Chapter 2 describes our first attempts to do large-scale distributed LOFAR processing on a shared infrastructure. We detail our successes with the LOFAR Radio Recombination Lines and Pre-Processing Pipelines, our software set-up as well as the limitations and future uses of this platform. This work is currently in prep.

Chapter 3 is our implementation for portable LOFAR processing on a massive scale. We show early results encapsulating LOFAR processing pipelines, and discuss future uses on other clusters. It is based on: **A.P. Mechev**, J.B.R. Oonk, et al. “An Automated Scalable Framework for Distributing Radio Astronomy Processing Across Clusters and Clouds”. In: *Proceedings of the International Symposium on Grids and Clouds (ISGC) 2017, held 5-10 March, 2017 at Academia Sinica, Taipei, Taiwan (ISGC2017)*. Online at <https://pos.sissa.it/cgi-bin/reader/conf.cgi?confid=293>, id.2. Mar. 2017, p. 2. arXiv: 1712.00312 [astro-ph.IM].

Chapter 4 discusses collecting and studying detailed performance statistics from automated LOFAR processing. It is based on: **A.P. Mechev**, A. Laat,

<sup>4</sup>[https://github.com/apmechev/GRID\\_LRT/](https://github.com/apmechev/GRID_LRT/)

<sup>5</sup>[https://github.com/apmechev/GRID\\_PiCaS\\_Launcher](https://github.com/apmechev/GRID_PiCaS_Launcher)

<sup>6</sup><https://github.com/apmechev/AGLOW>

et al. “Pipeline Collector: Gathering performance data for distributed astronomical pipelines”. In: *Astronomy and Computing* 24 (2018), pp. 117–128. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2018.06.005>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133718300490>.

Chapter 5 describes the capabilities of the initial workflow manager for automatic processing of LOFAR SKSP data. It shows several different scientific workflows and is based on: **A.P. Mechev**, J.B.R. Oonk, et al. “Fast and Reproducible LOFAR Workflows with AGLOW”. in: *2018 IEEE 14th International Conference on e-Science (e-Science)*. Oct. 2018, arXiv:1808.10735. DOI: 10.1109/eScience.2018.00029. arXiv: 1808.10735 [astro-ph.IM].

Chapter 6 describes a parametric model of resource usage for the LOFAR prefactor pipeline. We discuss difficulties that may arise from scaling LOFAR data, as well as the utility of our modelling method for SKA-size data. This chapter is based on **A.P. Mechev**, T.W. Shimwell, et al. “Scalability model for the LOFAR direction independent pipeline”. In: *Astronomy and Computing* 28 (2019), p. 100293. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2019.100293>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133719300290>.

Finally, Chapter 7 demonstrates the flexibility of our AGLOW software by implementing a full Continuous Integration pipeline for LOFAR software containers and LOFAR scientific pipelines. This pipeline can verify software pipelines against a test data set and can be used to track the data quality of data products as the processing software evolves. This work is submitted to *Astronomy and Computing*.



The contents of this chapter are based on a manuscript to be submitted to Astronomy and Computing

## 2.1 Introduction

In recent years, the need for multi-user oriented surveys supporting a large astronomical community with diverse scientific topics has increased, as this may maximize the scientific return value for a given amount of observation time. However, to encompass a large range in scientific topics within a single survey, it is often necessary to take the data at a common resolution (in terms of space, frequency and time) that supersedes the needs of the individual use cases. Such surveys therefore also require an increase in data transport, storage capabilities and (post-)processing power.

Here we will discuss the need for, and implementation of a large-scale compute platform to process data sets obtained with the Low Frequency Array (LOFAR). LOFAR is a modern radio telescope observing the radio sky at low frequencies from 10 to 240 MHz [20]. Its flexible observing setup means that it supports a large range in observing modes and settings that are utilized by a diverse user community with a broad range of scientific interests and goals. Some of the main science goals for LOFAR are established through the large key science projects [KSP; 20], but many other independent projects are also performed through open skies time observing.

We present the LOFAR distributed, shared processing (LOFAR-DSP) platform. This platform represents a first step towards facilitating radio astronomical processing on high-throughput compute infrastructures within the Netherlands and across Europe. The implementation of LOFAR-DSP has initially focused on processing interferometric imaging data for the LOFAR Surveys Key Science Project (SKSP). However, being a generic high-throughput data processing platform, it also pertains to a broader scientific audience and next generation Big Data experiments such as the square kilometre array (SKA)[39].

For the SKSP, the typical  $\sim 8$ -14 TB size (depending on the level of compression) of a single 8 hr data set and the  $\sim 50$  PB size for the full SKSP survey imply that the traditional way of processing, where radio astronomers use their own interactive recipes and facilities<sup>1</sup>, is neither efficient nor feasible. In this light, the value of well defined and robust pipelines that run on a scalable platform, delivering tractable, science-ready data products can not be overstated. This is even more so with future radio telescopes, such as the SKA, through which radio astronomy will enter the Exabyte-scale era.

As radio astronomy software development remains on a fast moving track [e.g. 19, 40], we have developed LOFAR-DSP as a light-weight solution that allows for easy integration with rapidly evolving pipelines. The technical setup chosen borrows many elements that have been developed for distributed Grid computing. The reasons for doing this are explained in more detail in Sections 2.2 and 2.3.

LOFAR-DSP was developed primarily by and for the SKSP and the radio recombination line (RRL) processing teams. However, the platform supports a larger variety of complex LOFAR processing pipelines that require large-scale, distributed, high-throughput compute infrastructures [e.g. 27–30, 41, 42]. For the SKSP case we will show how this platform, together with the GRID\_LRT processing framework [35] and the AGLOW LOFAR workflow orchestrator [37], has enabled the petabyte-scale processing of LOFAR data. Since 2017 more than 8 petabytes of SKSP data have been processed using LOFAR-DSP. For the recent SKSP LOFAR Two Metre Sky Survey (LoTSS) data release I [24], our solution directly contributed to 17 of the 26 accompanying papers.

This paper is structured as follows. In Sect. 2.2, we discuss the nature of radio astronomical observations and their intrinsic parallelism in the context of LOFAR SKSP interferometric data. In Sect 2.3, we introduce the Grid infrastructure in the Netherlands. The components of the LOFAR-DSP platform and its

<sup>1</sup>These facilities typically range from a laptop computer to small clusters comprised of a handful of large desktop machines.

implementation are described in Sect. 2.4. In Sect. 2.5, we show two pipeline examples that are built on top of LOFAR-DSP and in Sect. 2.6, we discuss the deployment of LOFAR-DSP across different infrastructures. We end with a discussion of the current platform and possible future improvements in Sect. 2.7 and present our conclusions in Sect. 2.8.

## 2.2 LOFAR observations and archive

The raw and/or reduced data for observed LOFAR projects is ingested by the LOFAR radio observatory in to the LOFAR long term archive (LTA). The LTA is a federated and distributed data archive that is hosted by three data centers: SURF-sara<sup>2</sup> in the Netherlands, Forschungszentrum Jülich<sup>3</sup> (FZJ) in Germany and the Poznań Super computing and Networking Center<sup>4</sup> (PSNC) in Poland.

The goal of the LOFAR SKSP team is to perform a tiered survey of the entire northern hemisphere aimed at imaging the low-frequency sky at unprecedented spatial resolution, depth, and frequency coverage. This survey serves the scientific goals of about 200 researchers across Europe and its details are described in Rottgering, Braun, et al. [43].

The first tier of the survey, carried out with the LOFAR high band antenna array (HBA), is described in Shimwell, Röttgering, et al. [41] and Shimwell, Tasse, et al. [24]. The SKSP HBA data is observed with 1 sec and 3 kHz resolution. This is subsequently pre-processed by the radio observatory flagging and averaging pipeline [44, 45] that performs a first round of radio frequency interference removal and then averages the data in frequency to 12.2 kHz. Given the broad range in SKSP science goals, requiring different processing strategies, the LOFAR radio observatory does not process the data beyond this initial pre-processing stage and ingests it in to the LTA.

### 2.2.1 Archived data sizes and transport

The archived LOFAR measurements are managed using dCache<sup>5</sup> as a front-end data storage manager and protected by a combination of access control lists (ACL)

<sup>2</sup><https://www.surf.nl>

<sup>3</sup><https://www.fz-juelich.de>

<sup>4</sup><https://www.man.poznan.pl>

<sup>5</sup><https://www.dcache.org/>

and Grid native X.509 based certificates. The measurements themselves are stored on tape backends and need to be moved to temporary disk storage prior to retrieval. This data staging is handled via a request through the LTA archive interface<sup>6</sup>. The staging service interacts with dCache, which enables seamless integration between disk and tape storage. After staging, the user can download the data via either the LOFAR download server or through a variety of Grid data transfer tools. The former provides URLs that can be resolved through HTTP for data retrieval. The latter use SRLs (Storage URLs) and TURLs (Transfer URLs) that are resolved via Srm (Storage Resource Manager) and gridftp for data retrieval<sup>7</sup>.

The size of a typical, archived LOFAR SKSP data set is in excess of 8 TB for a single 8 hr observation of a 5 degree by 5 degree area of the sky<sup>8</sup>. These data sizes are often too large to allow for efficient transport from the LTA sites to the compute facilities at the home institutes of LOFAR users, unless dedicated (and often costly) network connections are considered. Furthermore, significant storage space would be required at these institutes, especially when considering that the processing pipelines inflate the data by factors 2–3 during processing.

This data size and transport problem implies that a different solution has to be found in order to further process and reduce these large datasets before the data is served to the user. With this goal in mind in 2015 we created the LOFAR e-infra group<sup>9</sup> to develop a bulk processing solution enabling LOFAR processing at the LTA sites themselves. These sites also provide access for researchers to HTC and HPC compute facilities that have fast connections to the data storage systems holding the LOFAR data, thus eliminating the data transfer and storage issues.

## 2.2.2 Radio astronomical data & parallelization

Interferometric radio astronomical observations measure visibilities. These visibilities are samples of the Fourier transform of the sky and are represented by complex numbers consisting of measured phases and amplitudes as a function of frequency and time. The independent nature of each visibility measurement allows for a natural separation of the data along time and frequency. We use this intrinsic parallelization of the data to effectively spread (significant parts of) the processing of each

<sup>6</sup>The web interface is hosted at <https://lta.lofar.eu> In addition a python based application programming interface (API) also exists.

<sup>7</sup><https://www.dcache.org/manuals/Book-3.2/config/ch13s07-fhs.shtml>

<sup>8</sup>For related working groups such as the RRL processing group a single data set can be as large as 100 TB

<sup>9</sup><https://www.universiteitleiden.nl/en/research/research-facilities/science/lofar-e-infrastructure-group>

large data set over many small, independent jobs. This setup has been demonstrated in the framework presented in [35].

The embarrassingly parallel nature of the data processing, coupled with the large data sizes and I/O intensive nature of the processing means that HTC clusters are most suitable to carry out radio astronomical processing of LOFAR data. Furthermore, considering the fact that the archived data is stored in dCache managed Grid storage naturally leads us to consider a Grid computing solution for the LOFAR Surveys processing.

The Grid computing solution invoked here is based on elements that were originally designed for the Worldwide LHC Computing Grid (WLCG) project by the WLCG collaboration [46] and the various Grid initiatives. The Grid elements that are part of LOFAR-DSP are described in more detail in Sect. 2.4. The Grid connects a network of heterogeneous and distributed compute clusters to meet the massive compute requirements of I/O intensive data processing projects, such as the WLCG collaboration. Each compute node within and across Grid sites can be seen as an isolated island, i.e. it has contact to the outside world via internet and the workload manager, but it is not connected to the other compute nodes through an interconnect or a shared filesystem. This has advantages in that the use of a local file system and local scratch space is very efficient for I/O intensive jobs. However, this also has disadvantages in that the orchestration and movement of data, software and processing scripts becomes more elaborate, as compared to having access to the compute nodes over a shared file system.

Radio astronomical workflows differ from traditional high energy physics (HEP) jobs in that, in many cases, they are not completely parallel. Although, initially the SKSP workflows, on archived data, start out as highly parallel there are typically several tasks within these workflows that require all data to be brought together in order to derive better solutions and hence deeper imaging. Examples of such workflows are provided in e.g., [18], [35], [24] and [40]. In these, more complex workflows it is an advantage to also have the (large) intermediate data products, that need to be combined, collected at a site with high speed data transfer connections between the (temporary) storage back-end and the compute cluster.

## 2.3 LOFAR processing on distributed compute systems

The intrinsic data parallelism, I/O intensive nature of the workflow tasks and large data sizes involved for LOFAR all imply that we need a solution suitable for high

throughput compute that is connected over a fast network to the LTA. As discussed in Sect. 2.2.2, Grid computing offers this solution. In addition to the data size and transport issues, the LOFAR workflow tasks executed on the data can also be demanding in terms of memory and scratch space. For example, the SKSP direction dependent DDF pipeline discussed in [19] and [24] requires minimally 256 GB RAM and 3 TB of scratch space to complete. Contrary to this, the parallelized implementation of the SKSP direction independent prefactor pipeline discussed in [35] only requires a minimum of 8 GB RAM memory and 50 GB of scratch space for the most demanding tasks.

Having been optimised for less demanding HEP processing tasks not all Grid clusters are capable of handling the requirements for SKSP processing tasks. Here we will focus on the Grid infrastructure in the Netherlands that is able to satisfy these requirements. Other Grid clusters and facilities will be discussed in Sect. 2.7.

### 2.3.1 SURFsara & the Dutch Grid infrastructure

The Dutch National Grid Initiative (NGI) is a node of the European GRID initiative (EGI). It is hosted by SURFsara and Nikhef in Amsterdam. The Grid infrastructure provided at SURFsara is well suited for LOFAR data processing<sup>10</sup>. The Grid compute resources at SURFsara are provisioned on a per core basis to Grid jobs. This provisioning guarantees minimally 8 GB RAM memory and 80 GB scratch space per requested core. Compute nodes with up to 40 cores are available. In total the cluster provides 58 TB RAM memory, 2.3 PB scratch storage and 7400 cores.

The Grid cluster has a fast connection to the more than 60 dCache disk pool nodes that are also configured as doors and that serve the dCache managed data at SURFsara<sup>11</sup>. Each Grid compute node has a  $2 \times 25 \text{ Gbit s}^{-1}$  network connection and the total network bandwidth between the dCache managed Grid storage at SURFsara and the Grid cluster is  $1.2 \text{ Tbit s}^{-1}$ .

SURFsara also provides a dedicated user interface (UI) machine for Grid projects. This UI is aimed at easing the interaction between the Grid architecture and its users by having a rich set of Grid software and tooling pre-installed. The UI setup is identical to that of the worker nodes on the Grid cluster in terms of the supported software configuration. This enables the users to test and debug their processing pipelines on the UI before porting them to the Grid.

<sup>10</sup>[http://doc.grid.surfsara.nl/en/latest/Pages/Service/system\\_specs.html](http://doc.grid.surfsara.nl/en/latest/Pages/Service/system_specs.html)

<sup>11</sup>dCache currently manages over 50 petabytes of data at SURFsara.

LOFAR has a dedicated UI called `loui`, and LOFAR Grid jobs are submitted from `loui` to the SURFsara Grid cluster using the `gLite` middleware software. From `loui` it is also possible to submit Grid jobs to other Grid clusters; we will discuss this in Sect 2.7. Here we consider `loui` itself to be part of the SURFsara infrastructure (See in Figure 2.1). The LOFAR-DSP platform at SURFsara is deployed on `loui` (see Sect. 2.4).

## 2.4 LOFAR distributed shared processing platform

The processing framework (Grid\_LRT) and workflow orchestrator (AGLOW) for our LOFAR Grid processing solution have previously been presented in [35] and [37]. Here we present the underlying platform that we have named LOFAR-DSP. This platform re-uses and combines individual building blocks that were designed for high throughput GRID processing. LOFAR-DSP was first built in 2015 by the LOFAR e-infra group and has been continuously updated and refined since then.

Large experiments, such as LOFAR, that generate petabyte-sized datasets typically have lifetimes in excess of 10 years. However, compute technologies and their underlying software and hardware have a typical lifetime of maximally five years. This mismatch implies that portability becomes a very important aspect in the design of a long-term processing solution.

The overall aim of our LOFAR Grid\_LRT processing framework is to have as few dependencies as possible and thereby enable portability across Grid clusters and other compute facilities. The LOFAR-DSP platform provides the interface between this dedicated processing framework and the (generic) compute infrastructure. The platform itself is also portable, as we will discuss in Sect. 2.4.6 and 2.7.

LOFAR-DSP consists of 4 elements, the: (i) LOFAR software, (ii) workload manager, (iii) PiCas client, and (iv) Grid tools. These four elements are shown graphically in Figure 2.1. Here we will first discuss the requirements imposed upon the platform by the Grid\_LRT framework. Following this we will describe each of the four main elements that together comprise the platform, their connection and their interfaces.

### 2.4.1 LOFAR Grid\_LRT requirements

The LOFAR Grid processing framework, as presented in [35], has a few basic requirements. These are, (a) access to the LOFAR dCache managed Grid storage, (b)

a VOMS client, (c) access to the LOFAR software, (d) Python 2.7 or higher, and (e) outbound internet connectivity to connect the local job to the main job management database. This job management database is hosted on a CouchDB instance at SURFsara, accessed through an HTTP connection, and considered here to be part of the infrastructure.

The concept of Pilot jobs, see Sect. 2.4.5 and Figure 2.2, within the setup of the framework means that pilot job submission is part of the platform rather than the framework. Similarly pilot job scheduling is considered here to be part of the infrastructure. Job definition (e.g. input and tasks), orchestration and management however are handled by the Grid\_LRT framework<sup>12</sup> and the AGLOW<sup>13</sup> workflow orchestrator.

The LOFAR-DSP platform fulfils the requirements of the Grid\_LRT framework by providing the software tools necessary to access to the software, data, job scheduler and the job management database.

## 2.4.2 Grid tools

The LOFAR-DSP platform defines the interface between the data storage system and the processing framework. Archived LOFAR data is stored at the LTA sites (Sect. 2.2). For efficiency reasons we have chosen the Grid tools for data transfers in the Grid\_LRT framework.<sup>14</sup>

The Grid native data transfers tools are generic and are therefore considered to be part of the platform rather than the framework. The data transport tools selected as part of the LOFAR-DSP platform are `globus-url-copy`<sup>15</sup>, `uberftp`<sup>16</sup> and `GFAL2`<sup>17</sup>. The authorisation necessary to access LOFAR data requires a valid X.509 certificate and membership of the LOFAR virtual organisation (VO) in order to create an associated X.509 proxy. The proxy is created using the client `voms-client3` and the tools mentioned above will then use this proxy to authenticate with the Grid data storage system (i.e., via dCache).

<sup>12</sup>[https://github.com/apmevhev/GRID\\_LRT](https://github.com/apmevhev/GRID_LRT)

<sup>13</sup><https://github.com/apmevhev/AGLOW>

<sup>14</sup>The limited connectivity of the LOFAR download server means that the Grid data transfer tools provide data transfers that are on average more than order of magnitude faster when the network allows for it.

<sup>15</sup><https://www.globus.org/>

<sup>16</sup><https://github.com/JasonAlt/UberFTP>

<sup>17</sup><https://dmc.web.cern.ch/projects/gfal-2/home>

### 2.4.3 Workload management

The LOFAR-DSP platform defines the interface between the workload management system (or job scheduler) and the processing framework. At SURFsara, access to the Grid cluster is provisioned via a variety of middleware software. This middleware software does not directly schedule the jobs on a local Grid cluster. Instead, the middleware provides and translates the Grid job to a format that can be understood by the job scheduler (e.g. torque, pbs, Slurm) of the local Grid cluster. The LOFAR-DSP platform uses the gLite<sup>18</sup> middleware and hence contains this software to interact with the Grid workload management system. In order to submit jobs to a Grid cluster via gLite we need a valid X.509 proxy and the LOFAR VO needs to be registered on that cluster. The current gLite software is nearing its end of life and in Sect. 2.7 we will discuss possible alternative software.

The local job scheduler typically varies between different clusters and for non-Grid clusters there is no middleware layer to translate jobs to the required local format. Therefore this part of the LOFAR-DSP platform is less generic and can only be applied to Grid-published clusters. Fortunately, our use of Pilot jobs means that it is straightforward to change this part of the platform and accommodate different workload management systems and job schedulers. In Sect. 2.6 we provide an example where we run a modified version of the LOFAR-DSP platform on a cloud-based compute cluster with Slurm as the local job scheduler.

### 2.4.4 LOFAR Software

The LOFAR-DSP platform defines the interface between the LOFAR software distribution and the processing framework. The Grid resources, both locally and globally, are inherently heterogeneous and not accessible through a shared filesystem. In order to have a consistent LOFAR software stack available on all worker nodes we need a uniform way to distribute and compile the software.

To compile the LOFAR software we initially used the Softdrive<sup>19</sup> virtual drive solution offered by SURFsara. Softdrive offers a software environment that is identical to that of the SURFsara Grid cluster. Software compiled within this environment is therefore less likely to encounter errors upon execution locally at SURFsara.

---

<sup>18</sup><http://repository.egi.eu/>

<sup>19</sup>[http://doc.grid.surfsara.nl/en/latest/Pages/Advanced/grid\\_software.html](http://doc.grid.surfsara.nl/en/latest/Pages/Advanced/grid_software.html)

The compiled software on Softdrive is then distributed across the Grid worker nodes via the `softdrive.nl` directory as part of the CERN VM-Filesystem<sup>20</sup> (CVMFS). CVMFS is optimised to deliver software in a fast, scalable and reliable way. It is implemented as a POSIX read-only file system in user space. Files and directories in CVMFS are hosted on standard web servers and mounted in the universal namespace `/cvmfs`. For LOFAR-DSP, we host our software in `/cvmfs/softdrive.nl`. The `softdrive.nl` directory is linked to the Softdrive virtual drive and maintained by SURFsara.

CVMFS can be mounted on most computers and clusters. However, the software compiled within the Softdrive environment typically only applies to systems with a matching operating system and similar hardware. To deploy the LOFAR software across different infrastructures the Softdrive solution for compilation is therefore insufficient. Software containerization, as provided by for example Singularity and Docker, enables software to be abstracted from the environment in which they run. This allows us to port the LOFAR software across different compute systems. Since late 2017 we have containerized the LOFAR software using Singularity and provided software images<sup>21</sup> that we distribute via the `softdrive.nl` directory in CVMFS. We chose Singularity as, contrary to Docker, it enables us to execute the software image in user space.

Containerized software and the associated images provide an excellent first step in abstracting the LOFAR software from the local operating system and software environment. However, it does not fully abstract the LOFAR software from the underlying hardware. Important here are for example CPU instruction sets. If a software image is compiled on a system that has a CPU instruction set which is not compatible with that of the compute system where the image is executed then the software will very likely fail.<sup>22</sup> To eliminate this problem for the LOFAR software we have setup a KVM-based virtual machine (VM) that emulates the lowest common denominator in the accessible Grid hardware for LOFAR Grid jobs. This VM is used for LOFAR software compilation and hosted on the HPC Cloud<sup>23</sup> system at SURFsara.

We have tested the performance of common LOFAR processing tasks for both natively compiled and CVMFS-hosted software in [38]. In that work, we found no significant difference in performance between native compilation and soft-

<sup>20</sup><https://cernvm.cern.ch/portal/filesystem>

<sup>21</sup>Our full LOFAR Singularity images have sizes of 5–10 GB. Removing unnecessary source code and invoking squashfs compression we can reduce these images to sizes of 1–2 GB.

<sup>22</sup>Typically an error of the ‘illegal instruction’ is cast by the system.

<sup>23</sup>TBD - provide weblink for HPC cloud

|                 | image on SingularityHub | image on Softdrive    | image on Grid Storage     | CVMFS install |
|-----------------|-------------------------|-----------------------|---------------------------|---------------|
| Authentication  | none                    | none                  | grid proxy                | none          |
| Download time   | ~minutes                | instant               | ~seconds                  | instant       |
| Requirements    | singularity             | Singularity and cvmfs | singularity and gridtools | cvmfs         |
| Deployment time | instant                 | ~minutes              | instant                   | ~minutes      |

Table 2.1: Pros and cons of the different distribution methods for the LOFAR software. Deployment time refers to the time taken for the compiled software to be accessible at the processing nodes. The size of the software image is 1.3GB. Likewise, the entire CVMFS install is 1.7GB.

ware distributed by CVMFS. Similarly no significant differences in performance are found between natively compiled software and Singularity-based software images for LOFAR software.

Singularity images can also be compiled, hosted and versioned on SingularityHub. The downside of remote hosting of software images is the transfer time for those images, which can become comparable to the processing time for short jobs. We visualise the pros and cons of different distribution methods in Table 2.1.

### 2.4.5 Job management: PiCas & CouchDB

The LOFAR-DSP platform defines the interface between the job management database and the processing framework. The Grid\_LRT framework makes use of the PiCas pilot job workflow<sup>24</sup> and the LOFAR-DSP platform therefore includes the PiCas client.

The PiCas pilot job workflow was created by SURFsara as a light-weight Pilot job framework<sup>25</sup> that is easily adaptable and extendable. The central server for the PiCas framework is based on a web accessible CouchDB<sup>26</sup> database. For LOFAR-DSP this central job database is hosted by SURFsara and considered to be part of the underlying infrastructure, see Fig. 2.1.

Prior to pilot job submission, the central job database has to be populated. This is done by a set of dedicated PiCas CouchDB scripts that generate so-called job tokens<sup>27</sup> containing the required job input and tasks for a pilot job. Pilot jobs

<sup>24</sup>[http://doc.grid.surfsara.nl/en/latest/Pages/Practices/picas/picas\\_overview.html](http://doc.grid.surfsara.nl/en/latest/Pages/Practices/picas/picas_overview.html)

<sup>25</sup>[http://doc.grid.surfsara.nl/en/latest/Pages/Practices/pilot\\_jobs.html](http://doc.grid.surfsara.nl/en/latest/Pages/Practices/pilot_jobs.html)

<sup>26</sup><http://couchdb.apache.org/>

<sup>27</sup>These Tokens are CouchDB documents

submitted to a compute cluster are like regular jobs, but instead of executing a task directly they contact a central server once they are running on a worker node. Then, and only then, will they be assigned a task, retrieve data and start executing.

The central server handles all requests from pilot jobs and keeps a log of what tasks are running, are finished, and can still be handed out. This enables a powerful way of conducting central administration and management of jobs across a set of distributed compute resources. Only one central database is required to serve job input to pilot jobs running on any of the Grid worker nodes and across Grid clusters. Similarly, the same database is used to serve job input to pilot jobs running on any other type of compute cluster. Examples of LOFAR-DSP pilots jobs that run on cloud-based compute clusters and HPC systems are provided in Sect. 2.6

At SURFsara, LOFAR-DSP based pilot jobs are submitted via gLite to the Grid cluster. Once the job lands on a worker node it contacts the PiCas server for job input. During execution the status of a pilot job is tracked via the PiCas client and the associated job token is updated in real-time within the CouchDB database. The web frontend interface of PiCas also enables quick sorting of all job tokens into user defined views that provide real-time monitoring of all jobs within the database.

Light-weight pilot job frameworks, such as PiCas, excel at orchestrating very large numbers of independent pilot jobs. The generic PiCas framework is not intended to handle dependencies between individual pilot jobs in the queue that are considered to be tasks within a larger interconnected, complex workflow. To provide this higher level of orchestration we have built the AGLOW workflow orchestrator on top of PiCas [37].

#### 2.4.6 LOUI & SPUI

Together, the elements mentioned above comprise the LOFAR-DSP platform. Given that all these elements are software, they can easily be bundled and deployed via a VM or a container solution. At SURFsara, we have chosen to deploy LOFAR-DSP as part of the `loui` VM. This VM is also the login node for all LOFAR Grid users and uses a generic test environment for deploying new LOFAR workflows to the Grid.

The generic LOFAR user environment offered by `loui` is not optimal for operating continuous LOFAR processing in a stable and automated sense. This is due to the limited resources available on the VM and interference from other

users. To achieve continuous processing, a managed and dedicated environment is needed. This environment is provisioned via a second VM called `spui` (surveys processing user interface). This VM, managed by SURFsara and accessible to only the SKSP processing team, contains the stable and validated versions of the LOFAR-DSP platform. To enable continuous and automated LOFAR SKSP processing on `spui`, a robot proxy has been installed that automatically renews the validity of the X.509 proxy, and regular LOFAR-DSP pilot job submission is scheduled via `cron`.

The use of the latest validated LOFAR surveys workflows and processing software is not yet part of a fully developed and implemented continuous integration and continuous deployment (CI/CD) process. However, successful attempts, using Github, Travis, Jenkins and Singularity, are being explored by the processing team to first achieve continuous integration (Mechev et al. in prep.) and later also continuous deployment.

## 2.5 Executing LOFAR workflows on LOFAR-DSP

The LOFAR-DSP platform provides the foundational layer on which the LOFAR SKSP processing is run. The implementation of the LOFAR SKSP prefactor direction independent continuum calibration and imaging pipeline is discussed in [35] and [37].

The LOFAR-DSP platform can and is supporting a larger variety of LOFAR processing pipelines, e.g. pre-processing, spectroscopy, long baseline imaging and polarimetry for interferometric data, as well as spectral-imaging for tied-array data. As examples we will here briefly describe two pipelines that make use of the LOFAR-DSP platform; (i) direction independent spectral calibration (DISC), and (ii) the LOFAR Grid pre-processing pipeline (LGPPP). Both examples apply to interferometric data. The second case highlights our first approach towards abstracting the users from the compute and storage environment and enabling them to upload new jobs via interaction with PiCas only.

### 2.5.1 Direction Independent Spectroscopic Calibration – DISC

One of the goals of the LOFAR spectroscopy group is to use the SKSP measurements to discover and detect radio recombination lines (RRLs). The pipeline for direction independent spectroscopic calibration (DISC) was created to process LOFAR

data for spectroscopic studies, with a primary use-case of targeting bright, extra-galactic sources with the high-band antennas [28]. These observations are processed at a frequency resolution (3 – 12 kHz) that is 4–16 times higher than required by standard SKSP continuum processing, and additional steps to calibrate the bandpass are implemented during processing (Emig et al. submitted).

DISC uses the Grid\_LRT framework to define the interface between the processing workflow and the LOFAR-DSP platform. Parallel jobs are distributed in some steps to independently process the data efficiently. However, simultaneous, band-wide analysis of the full data set is also needed. In total, three steps of parallel processing the data are interwoven with two steps of band-wide processing. The intermediate products saved at each step together with the final products account for roughly 740 GB for a typical observation.

The high spectral resolution and detailed bandpass corrections needed in DISC processing require  $\sim 10^4$  CPU corehours for a typical SKSP data set. This is an order of magnitude more than required for `prefactor` direction independent continuum calibration [35, 38].

### 2.5.2 LOFAR Grid pre-processing pipeline – LGPPP

LOFAR-DSP supports a range of dedicated, complex processing pipelines for a variety of scientific goals. These pipelines have largely been created by specialist teams and ported to the Grid\_LRT framework by the LOFAR e-infra group. Although these pipelines often are publicly available, it can be daunting for non-specialists to acquire the necessary insights and skills to run these pipelines. We have therefore identified two other groups of users, (i) non-expert users affiliated with a specialist team for the required processing (often a KSP) and (ii) non-expert users with no affiliation to a specialist team.

The former group will typically be served by the specialist team with science ready products, as is the case for the SKSP. For the latter group such a specialist team is missing<sup>28</sup> and on-demand (re-)processing of LOFAR data via a non-expert interface is not readily available. This means that these users will not be able to obtain the science ready products that they need. To gain experience with how one may serve these non-expert users we have created the LOFAR Grid pre-processing pipeline (LGPPP).

---

<sup>28</sup>The LOFAR radio observatory team also offers some support for non-expert users, but this is naturally limited by the available resources

The LGPPP pipeline, built within the Grid\_LRT framework, represents a first step towards abstracting the user/researcher from the details of the processing and storage environment. It serves as an example of a pre-defined pipeline that a non-expert user can interact with by modifying only a few basic parameters. LGPPP is limited in scope to providing LOFAR users with the possibility to reject bad data points, reduce the LTA-stored data size and retrieve this reduced data. As a test case, LGPPP is implemented as a single New Default Pre-Processing Pipeline [NDPPP; 45] run, providing, in a predefined order, data flagging, averaging and demixing on a per Subband basis. In LGPPP the user can therefore only modify the step parameters for averaging and demixing<sup>29</sup>, decide on whether or not to carry out demixing, and provide a list of sources to be demixed. In addition, the user needs to provide as input a list of SURLS for the data sets to be reduced with LGPPP. This list is obtained by the user from the standard LTA interface.

An important requirement for LGPPP is that it must be able to run as a standalone service. Hence, we build on our existing LOFAR-DSP and Grid\_LRT solution and provide the user with the PiCas python client and a LGPPP job token generation script that takes as input the user defined parameters and the SURL list of datasets. Through a dedicated PiCas username and password the user is then able to populate the PiCas pilot job queue and monitor this queue via the CouchDB web interface. On `spui`, a running cron job activates the LGPPP pipeline that then checks whether there is work to do in the LGPPP queue and if so executes that work. The results of the LGPPP pipeline are shared with the user through an open WebDAV accessible storage managed by dCache.

The LGPPP pipeline and interface have successfully been tested with a selected group of non-expert users and it has been used as a first step towards achieving scientific results in [29]. The LGPPP approach can easily be extended to other, more complex LOFAR pipelines. However, as we will discuss in Sect. 2.7, we find that this is only useful for slow moving, robust pipelines. The fast moving targets of complex LOFAR pipelines imply that considerable effort is needed to maintain them for pre-defined execution and the LOFAR e-infra group decided to only pursue further implementation once this situation has become sufficiently stable.

---

<sup>29</sup>Demixing is a process in which contributions from bright off-axis sources to the observed measurements are estimated and removed [47]

## 2.6 Deploying LOFAR-DSP

Beyond the SURFsara Grid cluster a heterogeneous set of compute resources exist that could also be used for processing LOFAR data. Here we describe the deployment of LOFAR-DSP on a variety of systems, focusing on 4 cases: (i) federated Grid clusters, (ii) HPC systems, (iii) edge compute systems, and (iv) cloud computing.

### 2

#### 2.6.1 Federated Grid clusters – PSNC

PSNC is a partner in the LOFAR LTA. For reasons of data size and transport efficiency, discussed in Sect. 2.2, the LTA partners are natural candidates for processing LTA data. The Eagle HPC cluster at PSNC has been published as a Grid cluster and its worker nodes have outbound internet access. It is accessible for Grid jobs through the gLite middleware via a dedicated CreamCE queue. PSNC has also kindly installed the necessary Grid tools, CVMFS and Singularity on Eagle.

It is therefore not necessary to port the LOFAR-DSP solution to a federated Grid site such as PSNC. Instead, we can submit our LOFAR Grid jobs to Eagle directly from `loui` in the same manner as for the SURFsara Grid cluster. The only difference here being that we need to select the appropriate CreamCE. This functionality, of course, is the very essence of Grid computing.

There is one other important difference between the SURFsara Grid cluster and the Eagle cluster. The Eagle cluster, being an HPC system, has very little scratch space on a worker node for local processing. Instead, Eagle has a shared file system with associated globally mounted home and project directories. It was therefore necessary to adjust the LOFAR workflow scripts to take the local storage difference into account.

This change in local storage was straightforward to implement in `Grid_LRT` and performed by re-mapping the `TMPDIR` environment variable that is used for creating a unique, temporary processing space for each of our jobs. The change is reflected in the PSNC branch of `Grid_LRT` scripts repository on Github. LOFAR SKSP Grid processing jobs on Eagle were successfully tested in 2017 and 2018.

### 2.6.2 HPC systems – FZJ

FZJ, via the German GLOW consortium, is also a partner in the LOFAR LTA. FZJ hosts a variety of compute systems. For LOFAR SKSP data the JUWELS HPC system is the most suitable. JUWELS is connected to the FZJ storage systems hosting the LOFAR data via JUDAC (Juelich Data Access Server).

Similarly to the Eagle HPC system, JUWELS relies on a shared file system for data processing. However, there are other important differences in that the JUWELS cluster is not published as a Grid cluster and hence is not accessible from `loui` via `gLite`. Furthermore, JUWELS does not support Singularity or CVMFS and its worker nodes do not have outbound internet access.

Given that the JUWELS can not comply with our basic `Grid_LRT` requirements, see Sect. 2.4, we decided not to port the LOFAR-DSP platform to JUWELS. Instead, we have developed and need to maintain a separate set of tools. These tools re-use and rely on some of the elements (e.g., `PiCas`, `Softdrive`) that are used in LOFAR-DSP and where possible we have aligned our JUWELS implementation with LOFAR-DSP.

LOFAR processing on JUWELS (and previously JURECA) is realised in two steps: (i) Software installation. We use a pre-compiled, non-containerized, LOFAR installation hosted on `Softdrive` at `SURFsara`, using the `parrot-connector`. Regular updates can be performed via a simple data transfer (e.g., `rsync`). (ii) Job management and execution. We developed a monitoring script,<sup>30</sup> continuously running on JUDAC, which acts as an interface between the `PiCas` pilot job database hosted at `SURFsara`, the LOFAR LTA in Jülich and JUWELS. Upon completion, the processing jobs at FZJ send the results to the dedicated SKSP Grid storage at `SURFsara` for further processing and distribution.

### 2.6.3 Edge computing

In addition to the large processing facilities at the LTA sites, the LOFAR-DSP platform can perform LOFAR processing on compute facilities located at research institutes and universities. The usefulness of these edge resources for orchestrated, large-scale processing depends on the available compute and storage resources, the network connectivity to the LTA sites and the level of IT support. Typically these

---

<sup>30</sup>SKSP\_monitoring.py

resources are local workstations or batch processing clusters that have not been published as Grid clusters.

By exchanging the Grid workload manager in LOFAR-DSP with the local job scheduler the platform can run on these local facilities. This makes it possible to further scale the scientific processing for LOFAR data. We successfully tested running Grid\_LRT prefactor jobs, using LOFAR-DSP, at the Leiden Observatory LOFAR cluster and the Herts cluster Hertfortshire. However, we found that these facilities at the edge are fundamentally limited by the bandwidth of their connection to the LTA.

This network limitation prevents large-scale SKSP processing of LOFAR data at these sites. Instead these sites are primarily used for performing development studies and high-level, post-processing of calibrated datasets obtained from LOFAR-DSP processing at the LTA sites. These, primarily direction independent calibrated data, are hosted on the SKSP Grid storage at SURFsara and have typically been reduced in size by a factor 16. Following this reduction, these data sets are much more easily distributed to the edge sites than the original LTA stored data.

## 2.6.4 Clouds

In the past decade, many public and private IT providers have begun offering on-demand cloud-based VMs for hosting a variety of services, such as web applications and lightweight data analysis. Cloud users can define the resource requirements of these VMs, and once launched, they have root access to the VM. This flexibility makes it possible for a user to tailor the VM to their needs. At the same time, this type of self-service mode of operations comes at the expense of significantly increased investment in knowledge, time and expertise from the user to exert the necessary control to keep their VM's running, updated and secure.

There are a number of reasons as to why cloud computing has not gained traction within the more data-intensive sciences such as radio astronomy. In addition to the expertise issues mentioned above, we can mention that: (i) radio astronomical workflows require detailed orchestration of complex pipelines. In the cloud environment, the VMs also require orchestration and thus become an additional management burden. (ii) Radio astronomy is pushing the CPU, memory and local storage limits of individual compute nodes. As such, it is a costly and time consuming process to get the necessary VMs launched within current cloud environments. (iii) Radio astronomy is pushing the limits of data throughput and persistent,

long-term storage. The required network and storage resources to run, for example, LOFAR SKSP processing can hence become very costly, especially upon considering that each data set will be processed several times by independent science teams with different versions of their pipelines.

The LOFAR-DSP platform can treat cloud VMs as processing resources, as long as the VM instance adheres to requirements listed in Sect. 2.4. The most straightforward way would be to set the VMs up as a batch processing cluster, install the required software and publish it as a Grid cluster. In the Helix Nebula Science Cloud<sup>31</sup> project we, together with experts from T-Systems and Rhea, managed to setup a Slurm<sup>32</sup> based batch processing clusters in the T-Systems and Rhea clouds. Although we did not publish these clusters as Grid clusters we were able to setup all other parts of the LOFAR-DSP platform and successfully carry out functional tests for Grid\_LRT based LOFAR processing. However, given the above mentioned investments and limitations we decided not to pursue large-scale cloud processing beyond these initial tests.

## 2.7 Discussion

The LOFAR-DSP platform and the Grid\_LRT framework, developed by the LOFAR e-infra group and SURFsara, have now been operational for almost 4 years. During this period we have processed over 8 PB of LOFAR data and we have experimented with extending both the framework and the platform to include more pipelines and additional compute facilities.

On average about 1.5 FTE per year has been dedicated to implementing, maintaining and further developing this processing project. This limitation in available human resources has made it difficult to go beyond the current implementation described here and make progress towards e.g., a full CI/CD implementation. Below we will briefly discuss some of the challenges that we are facing in the next couple of years. In particular, we will focus on the future of Grid computing and on reaching the necessary level of robustness and automation for the LOFAR pipelines that we need to process petabyte datasets in a semi-continuous manner.

<sup>31</sup>More information about this project can be found in the associated deliverables: D6.2 Integrating commercial cloud services into the European Open Science Cloud: <https://zenodo.org/record/2598039#.XVurI3vRYuU>, and D6.3 Demonstration to the EC of the test products resulting from the procured R&D services: <https://zenodo.org/record/2598060#.XVux73vRYuU>

<sup>32</sup><https://slurm.schedmd.com>

### 2.7.1 Future of our Grid computing for Astronomy

In Sect. 2.2 we discussed that radio astronomical data processing, from nearly raw data up to and including imaging, is very well suited to high throughput Grid processing. In assembling the LOFAR-DSP platform we have gratefully made use of existing Grid tools. Some of these tools are now nearing their end of life, in terms of support, and the platform will need to be updated accordingly in coming years (2020–2021). In particular we mention here that gLite-wms, gLite-ce, CreamCE, and globus-url-copy will need to be replaced.

The gLite workload management tools we plan to replace with DIRAC<sup>33</sup>. For pilot job submission of Grid\_LRT-based pipelines this change will be transparent in that the glite-wms-job-submit commands will be replaced by the corresponding dirac-wms-job-submit commands. For AGLOW, our workflow orchestrator, this change has a more pronounced effect in that the current implementation of AGLOW monitors the status of a Grid job via glite-wms-job-status and it parses the output retrieved from this command. The change in the output retrieved from dirac-wms-job-status, as compared the glite-wms-job-status, implies that the parser needs to be updated. Alternatively, it may be considered to move the AGLOW job monitoring from the workload management level to the (PiCas) job token level. This would allow for generic AGLOW monitoring in a manner independent of the different workload management tools and job schedulers. However, in this case we would also need to handle interrupts between the workload management tools and PiCas to ensure that all job tokens are always eventually updated to the correct state of the corresponding pilot job.

The replacement of the CreamCE, by e.g., ARC-CE or HTCCondor-CE, will be carried out at the infrastructure level. For the LOFAR-DSP platform and Grid\_LRT framework this should only amount to a change in the queue names provided to dirac-wms. For replacing globus-url-copy and the underlying gridftp protocol there are several options, but we expect to use gfal2 with either WebDAV or XROOTD as the underlying protocol. For dCache stored data in particular the combination of dCache macaroons with WebDAV is attractive as this would remove the dependence on X.509 certificates and enable the use of other more generic data transfer tools such as rclone and curl. Within the European Open Science Cloud (EOSC) hub project we are investigating the use of dCache macaroons for not only the SKSP, but also the wider LOFAR community.

---

<sup>33</sup><https://dirac.readthedocs.io>

### 2.7.2 Further automation and optimisation

Beyond the future of the Grid computing tools there are a number of other relevant developments that could be considered in further improving the LOFAR-DSP platform, as well as the framework and pipelines built on top. We will briefly discuss some of them here.

The dependence of the LOFAR-DSP framework on the Grid workload management system is both a strong and a weak point. It is strong in that the platform uses an accepted and powerful community standard to submit and distribute its processing jobs. It is weak in that this approach will only work for Grid clusters on which the LOFAR VO has been enabled and for which compute time has been granted. We show in Sect. 2.5 that it is straightforward to change the workload management system, to e.g., Slurm, but this has not yet been automated and would also require changes in AGLOW (Sect. 2.7.1).

SKSP processing is currently driven by a database that stores the processing state of all SKSP observations and which is regularly updated by querying the LTA for newly archived SKSP data. Instead of this poll-and-pull mechanism, there has been some development within the storage and processing communities to evaluate event driven processing. In the latter case, data arriving in the LTA may self-generate their processing chain via e.g., LTA ingest events or dCache events. For large archiving projects, where data is swiftly moved from disk to tape, such event driven processes may also ensure that the tape to disk staging latency is avoided.

X.509 certificates and their derived proxies enable seamless access to Grid storage and compute. However, they are not heavily used outside of the Grid world. In recent years the Grid community has been exploring token-based Authentication & Authorisation Infrastructure (AAI) solutions to replace X.509. If achieved, such solutions could also be extended to include other services upon which LOFAR-DSP is reliant, e.g. the PiCas CouchDB server and the LOFAR LTA interface. This may not only lower the administrative burden for LOFAR processing services, but also improve the user experience and lower the threshold for new users.

Although some initial progress is being made in generating a well defined chain for the maintaining, validating and deploying the platform, framework and pipelines, this chain is far from ready. We consider it highly beneficial for the LOFAR (SKSP) community at large if such a 3-phased CI/CD process could be implemented and automated. Such an implementation requires organisation and human resources beyond what is made available now. A dedicated team is not only

important during the initiation phases, but also a requirement for managing this process in the long term after the implementation phase. During the four to five years that the LOFAR e-infra group has been operational, this lack of human resources has been our main bottleneck in moving beyond the current implementation.

Finally we mention that the use of containers and software images have been a major step forward in our efficiency and ability to port the complex set of LOFAR software across differently managed infrastructures. One of the future goals of the LOFAR e-infra group is to offer all of our software in an as easy as possible executable format. This would enable a more unified processing environment and better reproducibility of LOFAR data processing. As such we foresee that we will also offer LOFAR-DSP as a software container. One important issue here is that not all IT providers have yet embraced containers, as these are sometimes seen as a security risk. However, for large, complex computing projects such as LOFAR SKSP it is becoming increasingly clear that we can no longer afford to continually optimise and update our software for different infrastructures with different operating systems and system libraries. This is even more true in the coming years where first LOFAR 2.0 and later the SKA will come online, and provide another increase in data rate and size.

## 2.8 Conclusions

The LOFAR radio telescope archives tens of terabytes of data daily, and its LTA grows by  $\sim 7$  PB per year. These rates are too high to transfer and process data at compute facilities that do not have a dedicated network connection to the LTA sites. To enable efficient processing of the archived LOFAR data for the SKSP project and dissemination of the results, we need to process the incoming LOFAR data with low latency and serve science-ready data sets and products to the astronomical community.

Here we have presented the LOFAR-DSP platform, which we show is a major building block towards enabling massively distributed LOFAR processing. We describe the underlying technology and our implementation of the platform. We note the strengths of our solution and provide suggestions for future improvements to make the LOFAR-DSP platform easier to use and scale across multiple infrastructure sites. The main contributions of the work presented here are the following:

- We present LOFAR-DSP as a platform for distributed processing of LOFAR data on a shared IT infrastructure. We focus on the Grid implementation at SURFsara and show examples of deployments on other IT infrastructures.
- We show the implementation of the PiCas pilot job framework for LOFAR processing jobs. We highlight how this framework is used as a central resource to distribute and monitor LOFAR processing jobs across different IT infrastructures.
- We provide two pipeline processing examples: (i) DISC for spectroscopic data processing and (ii) LGPPP for pre-processing. We show how these pipelines interface with LOFAR-DSP. For LGPPP we also show a possible path towards abstracting the non-expert user from the details of the underlying IT infrastructure.
- We discuss the process of assembling LOFAR-DSP, the underpinning (Grid) software life-cycle and provide suggestions for further improvements to the platform and the processing framework built on top.

The impact of LOFAR-DSP, in combination with the Grid\_LRT framework [35] and the AGLOW [37] workflow orchestrator, is evident from the recent data releases by the SKSP project [24, 41].

To date the LOFAR-DSP platform has processed over 1000 data sets for the LOFAR Two-Meter Sky Survey [LoTSS; 24] and a variety of other LOFAR projects, such as the RRL spectroscopic surveys [27, 28, 30]. The platform has contributed to the data reduction for more than 40 scientific publications. The scientific output of the SKSP surveys has been rapidly growing in recent years and this can be understood through: (i) the improved understanding of calibration and imaging techniques for LOFAR [e.g. 16, 18, 19, 44, 45], and (ii) the automation and massive scaling of the LOFAR pipelines [this work; 35, 37, 38].

These achievements indicate that the community is starting to understand what is needed for the upcoming challenges that will be posed by the SKA in the next decade and that will deliver data at a rate more than two orders of magnitude greater than LOFAR. From our experience with LOFAR, and soon LOFAR 2.0, it is clear that we need to continue to increase our investment in IT knowledge and

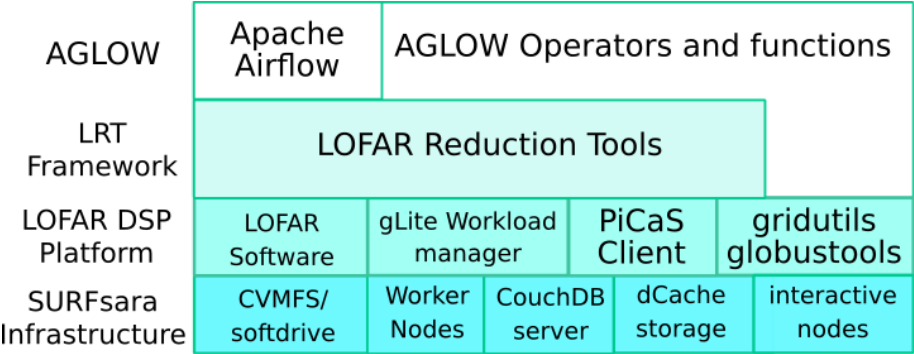


Figure 2.1: Structure of the LOFAR-DSP platform as it relates to the rest of the services used to process LOFAR data on the Dutch grid. Both the LOFAR Reduction Tools and AGLOW make direct use of components provided by LOFAR-DSP. Likewise, the LOFAR-DSP platform interfaces with the infrastructure provided by SURFsara.

infrastructure in order to support the future of radio astronomy and to strengthen the bonds between two scientific domains.

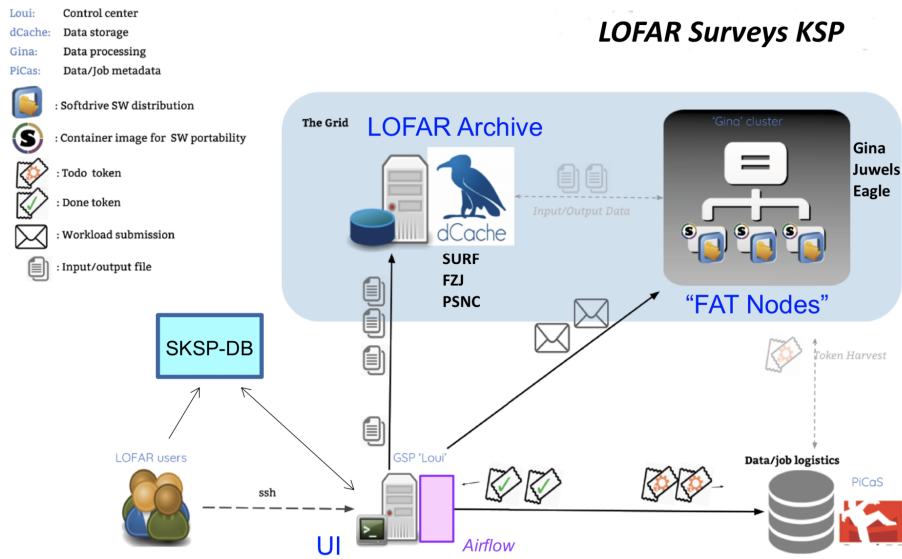


Figure 2.2: Implementation of pilot job launching by the LOFAR-DSP platform. LOFAR users log onto the UI machine and define jobs on the PiCaS database. They submit jobs through the grid middleware, and each job harvests an unprocessed PiCaS token, and executes the job defined therein. Higher-level orchestration is done by Apache Airflow running on the UI machine.



**Original Abstract:**

The Low Frequency Array (LOFAR) radio telescope stationed near Exloo, the Netherlands is an international aperture synthesis radio telescope used to study the universe at low frequencies. Aperture synthesis requires large amounts of computation between data acquisition and science ready images. The LOFAR Two Meter Sky Survey (LoTSS) will require to process 50 PB of data within five years. The data rates demanded by this project require processing at locations with high-speed access to the data. The current software packages are not suited for all cluster architectures, and cannot launch and monitor processing at multiple locations.

To complete the LoTSS project, the processing software needs to be made portable and moved to clusters capable of handling the data rates above. This work presents a framework that makes the LOFAR software portable, and is used to scale out LOFAR data reduction. The high throughput achieved will make imaging 3000 observations possible within five years.

This chapter is based on the work of **A.P. Mechev**, J.B.R. Oonk, et al. “An Automated Scalable Framework for Distributing Radio Astronomy Processing Across Clusters and Clouds”. In: *Proceedings of the International Symposium on Grids and Clouds (ISGC) 2017, held 5-10 March, 2017 at Academia Sinica, Taipei, Taiwan (ISGC2017)*. Online at <https://pos.sissa.it/cgi-bin/reader/conf.cgi?confid=293>, id.2. Mar. 2017, p. 2. arXiv: 1712.00312 [astro-ph.IM]

## 3.1 Introduction

The LOFAR radio telescope is the world's largest aperture synthesis array with more than 20,000 antennas and baselines of 60 m to 1000 km[20]. With its unprecedented sensitivity and angular resolution at ultra-low frequencies, LOFAR's goals are far-reaching: from studying pulsars and supernova remnants to the evolution of distant galaxies. Additionally, LOFAR is a pathfinder for the larger Square Kilometer Array (SKA) radio telescope. The SKA-Low telescope is expected to increase data size[48] to 400TB per day, creating more than 120PB per year[49].

The LOFAR Two Meter Sky Survey (LoTSS) Survey[41] will observe 3000 different fields that collectively will map the entire northern radio sky. The size of each of these data sets is 16TB, making for a total of 48 Petabytes. To complete the LoTSS survey in the project's anticipated 5 year duration, 1PB of data need to be processed each month.

To mitigate delays caused by data transfer, the reduction must be done at a location with a high bandwidth connection to the raw data. Software packages for the initial processing of LOFAR data already exist[18], however they were not designed to work on all cluster architectures. To complete the LoTSS project in time, a framework is needed to automatically process multiple data sets at a cluster with a fast connection to the data.

Building on previous work [50], we created a framework to launch and monitor processing for multiple data sets. We present this framework, built to process data sets across multiple machines. The framework is named the LOFAR Reduction Tools (LRT), and it provides:

- Automation, enabling processing of multiple concurrent jobs
- Portability, enabling processing at different locations
- Scalability, enabling adding worker machines as required by the workload
- Generalization, enabling the integration of software from other scientific domains

The LOFAR data reduction software known as 'pre-FACTOR'[40] was integrated in this framework. This software has been in use since November 2016 and at the time of writing (Feb 2017) had processed more than 100 data sets. This equals one-quarter of all the LoTSS data gathered from September 2014 to March 2017. By deploying the LRT framework on a cluster with a high-bandwidth connection to the data, the entirety of the LoTSS data can be reduced within the five year time span of the project.

The paper is structured as follows: Section 3.3 outlines the LOFAR data reduction process and computational requirements. Section 3.4 describes the design of the LRT framework and its capabilities, Section 3.4.2 describes the modification of the existing LOFAR software, and the performance and results are in Section 3.4.4. Finally, conclusions and future work are in Section 3.5.

## 3.2 Related Work

Scientific processing is increasingly moving to Grid and cloud-based distributed computing. With increasing data sizes, researchers have begun focusing on scalable ways to parallelize their workflows. From genetic sequencing [51] and bioinformatics[52] to neuroscience[53] and ecology [54], ever growing data sets have driven the development of distributed workflow systems in science[55][56].

The framework presented in this publication is built on previous work distributing LOFAR pre-processing on a computing cluster (Oonk+ in prep) using a PiCaS server to track progress[57]. The required infrastructure, a PiCaS[57] server and a CernVM Filesystem client[58] have already been deployed on the target cluster and tested prior to this work. Additionally, continual support for these packages was provided by the SURFsara science support group.

PiCaS[59] and CouchDB[60] have been used in other distributed computing projects to launch and monitor jobs and exchange metadata. Job monitoring using PiCaS is also used in projects such as Sim-City[61] and Finite Element modelling for sea dyke design[62]. In these works, pilot jobs were automatically launched and tracked remotely. The PiCaS framework enabled a high degree of automation and has helped process large amounts of data.

CouchDB is also successfully used by the LHCb team to monitor the nightly build process of their software[63] and by Sante et al.[64] to launch asynchronous jobs to visualize and analyze gene sequencing data. As CouchDB documents can hold arbitrary information and attachments, the use of the CouchDB platform favours projects requiring the storing of metadata for many concurrent jobs.

CernVM-FS[58] has been used by projects to package and publish software. Software such as the ATLAS [65] and the NO $\nu$ A [66] are compiled on a central server and published to worker nodes. The LOFAR software has been similarly packaged (Oonk+ in prep). This makes deployment of processing scripts possible without compilation on the worker machines.

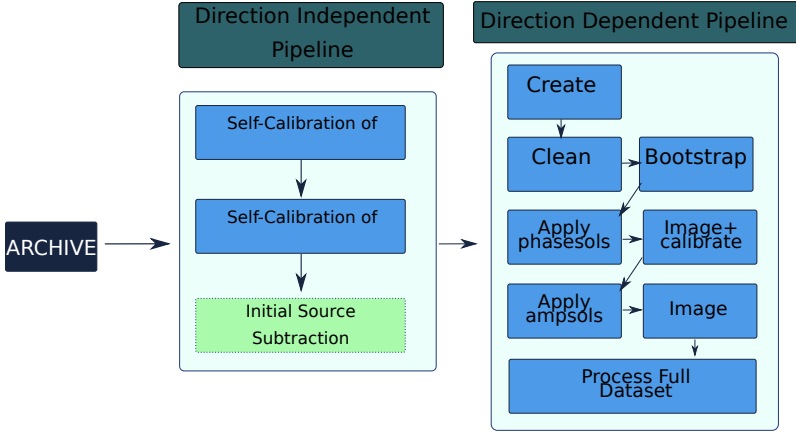


Figure 3.1: A schematic of the data flow through the Direction Independent Pipeline (pre-FACTOR[40]) and one of the Direction Dependent Pipelines (DDFacet). The Initial Source Subtraction is only necessary for some DD pipelines and is not currently implemented.

### 3.3 LOFAR Data Processing

Creating images from data collected by an aperture synthesis array[67] requires several steps of calibration and imaging. In order to place this work in the proper radio astronomy context, a brief introduction to LOFAR data processing follows. Section 3.3.1 gives an overview of LOFAR processing from an archived observation to a final image. A schematic of this is in Fig.3.1. Section 3.3.2 details the processing steps currently implemented as well as their computational challenges. Section 3.3.3 contains an overview of the benefits of integrating the processing software with the LRT framework. Finally, section 3.3.4 gives a description of the processing by focusing on the data flow. A visualization of the processing detailing the flow of data is presented in Fig. 3.2.

#### 3.3.1 Producing Images From LOFAR Data

The raw data is stored in the LOFAR Long Term Archive. Typically, an 8 hour observation results in a 16 TB data set split into 244 65GB files. Throughout the LoTSS data processing, this data is reduced to a 500GB set of calibrated data. The calibrated data is then imaged, producing a final set of a few 1.2GB images of 25k\*25k pixels. The calibrated set is archived as it can be re-imaged in the future.

Data reduction for the LoTSS survey is split into two pipelines: Direction Independent (DI) calibration and Direction Dependent (DD) calibration (Fig3.1). The Direction Independent pipeline[40][18] produces images that are limited in resolution and contain instrumental effects[25]. To achieve high fidelity continuum images, the ionospheric and beam errors must be corrected[68]. As those effects vary across the field of view, a Direction Dependent calibration step must follow. Without these corrections, the image quality and resolution is severely limited[18]. We are currently working on a GRID implementation of the Direction Dependent pipeline and will report on this in a future paper.

### 3.3.2 Direction Independent Processing

The LOFAR telescope consists of many antennae, each with its own electronic gain. A gain calibration needs to be performed by observing a bright calibrator source before or after the science target[68]. Using this observation, the antenna gains for the telescope are calculated. This is performed by the Calibration pipeline of the pre-FACTOR software[40]. The results from this step are applied to the science target, and target observation is averaged and processed. This consists of removing Radio Frequency Interference and subtraction of bright off-axis sources, and finally, calibration against a sky model derived from surveys conducted with previous telescopes[68][18]. These steps are performed by the Target pipeline of the pre-FACTOR software[40]. The result is a calibrated data set which is up to 64 times smaller than the uncalibrated archived data.

A 16TB data set cannot fit into a machine's memory, which is typically less than 128GB. Because of this, the Direction Independent processing is split by dividing the original full-bandwidth observation into independent chunks of narrower bandwidth. A typical observation spanning 48 MHz is split into 244 files called Subbands, each of which spans 0.1953 MHz. These Subbands are processed simultaneously, as each undergoes the same processing steps. This is a form of data-level parallelism.

The Direction Independent calibration pipeline (Fig.3.2) consists of an existing set of scripts which use the LOFAR software suite[69] to process the initial data sets. These scripts also handle the post-processing and application of the calibration results[18]. These scripts are contained in the package *pre-FACTOR*[40]. The order of the scripts and their parameters are contained in a parameter-set file (henceforth 'parset'). The parset defines a sequence of procedures, each launching

one or more executables. The execution of the procedures in a parset defines a step in the Direction Independent Pipeline.

### 3.3.3 Implementing the DI Calibration on the GRID

The LoTSS survey is mainly conducted by a team at Leiden University. The bandwidth between a University cluster and the LOFAR data archive is too low to download the data, 10 MB/s in the case of Leiden University. At Leiden University, downloading of one data set would take ten times longer than the processing. The processing was moved to the SURFsara grid location at the Amsterdam Science Park<sup>1</sup> as there were previous successes in processing LOFAR data at SURFsara (Oonk+ in prep). In order to take advantage of the computational resources at the SURFsara Gina cluster[70], the LOFAR pre-FACTOR pipeline was modified as part of the development LRT framework. The two steps of the DI reduction, the Calibrator and Target, were each split in two parts. The first part of the Calibrator and Target processing is parallelized by running one file per node, and the second runs combine these results. This takes advantage of the data level parallelism of LOFAR processing.

Additionally, splitting the computation makes it more robust. In the case that the download or processing of one job fails, it can be restarted without disrupting parallel jobs. When a step has finished processing, the next step can be launched automatically enabling the massive processing of LOFAR Surveys data.

The pre-FACTOR software was designed to be run on single machines or clusters with a shared file system. Because the worker machines at the SURFsara cluster have isolated storage, scripts are included in the LOFAR Reduction Tools to load the relevant data on the worker node before processing. After a job is finished, the scripts save intermediate results to a storage location external to the cluster.

### 3.3.4 Data Flow and Processing Steps

A researcher interested in processing LOFAR data needs to download it from the LOFAR Long Term Archive. Leiden University leads the LoTSS survey and has a computer cluster dedicated for LOFAR processing. The connection between the Leiden location and the LOFAR data archive is typically 10 MB/s. At this rate,

<sup>1</sup>[http://docs.surfsaralabs.nl/projects/grid/en/latest/Pages/Service/system\\_specs.html](http://docs.surfsaralabs.nl/projects/grid/en/latest/Pages/Service/system_specs.html)

downloading a single 16TB LoTSS data set completes in two weeks. At this rate, transferring 3000 data sets would take over a century.

Unlike at Leiden University, the SURFsara clusters have a gigabit connection to the data archive, accelerating the data retrieval to 1.5 days. Additionally, having two orders of magnitude more processing nodes than at Leiden, the reduction can be further parallelized. This has accelerated the processing of one (downloaded) data set from more than two days to less than half a day.

Using intermediate storage to hold the results from each step, the pre-FACTOR DI pipeline (Fig.3.1) was split into four steps as in Fig. 3.2. The Calib 1 and Target 1 steps download the raw data at one piece per worker machine and store the processing results (Calibration Tables and Processed data respectively) in storage. The second Calibration and Target steps combine these results and process them producing the calibration solutions and data sets respectively.

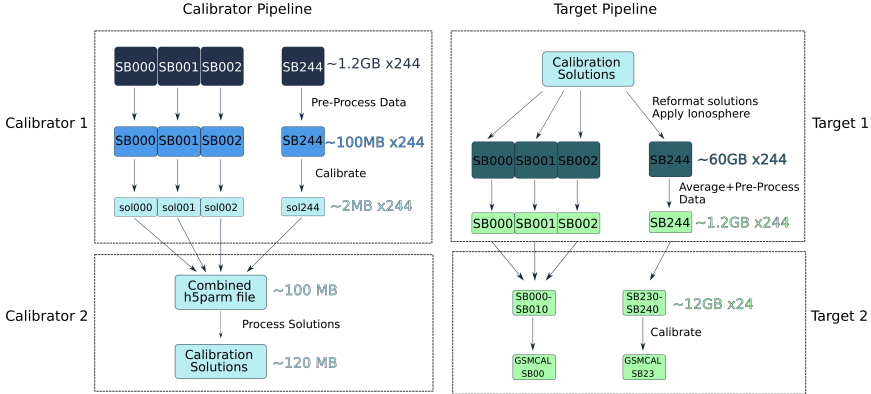


Figure 3.2: Data flow and parallelization of the Direction Independent Processing. The Calibrator 1 and Target 1 steps run concurrently as independent jobs. Calibrator 2 and Target 2 combine these results. Note that the Target 1 step requires the solutions produced by the Calibrator 2 step. This places a strict ordering on the processing steps.

### 3.4 Framework Design

The LRT framework (Fig. 3.3) was developed to automate the LOFAR Direction Independent calibration by processing the data at the Gina cluster SURFsara[70]. The goal of the framework was to adapt the pre-FACTOR package to take advantage of the large computational resources and high bandwidth at this site. Thanks to the data-level parallelism, using the computational resources at the Gina cluster accelerates the data reduction.

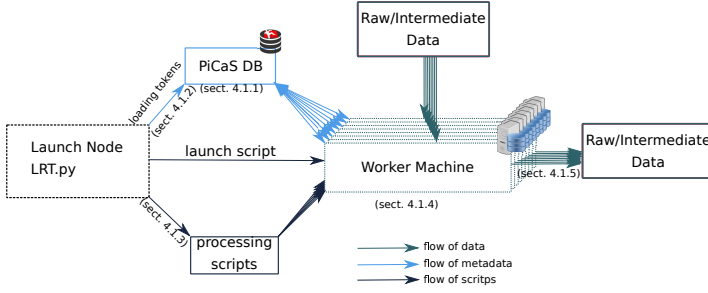


Figure 3.3: Overview of the design of the LRT framework.

### 3.4.1 Framework Elements

The LRT framework consists of a set of modules responsible for different parts of the data reduction. The `srmlist` module handles the links to the data. If the data is on tape, it sends a command to stage it to disk. The `sandbox` module creates an archive of processing scripts and uploads it to storage. The `Token` module is responsible for managing metadata, which defines a processing job. Appendix 3.A contains information on the functionality of each module and their use.

#### Storing Job Metadata

The LRT framework is effective for pipelines that execute the same processing steps on a large data set split across many machines. Each part of the data set is processed on a single machine, and the metadata of this job is stored in a remote database which can be read from and written to by the worker machine. By using a concurrent document-oriented database such as CouchDB[60], each document can store the metadata regarding a single processing job. This is not possible with relational databases such as MySQL. These documents are called Tokens, as defined by the PiCaS framework[57].

The first implementation of PiCaS and CouchDB for LOFAR data reduction was carried by J. B. R. Oonk and N. Danezi in the context of the LOFAR spectroscopy project and custom user processing (Oonk et al. in prep). This first implementation focused on processing individual data sets and required a high level of user interaction. Here we extend this implementation to automatically handle and connect multiple runs (calibrator and target) and their products.

By logging the name of the pipeline step in the job description, each job is aware of which pipeline step it is processing and executes the appropriate scripts. Important to note is that the CouchDB documents can store text and integer values as well as file attachments. The LOFAR implementation uses attachments to store diagnostic files produced by the worker machines, lists of links to the data and parset files that define the pre-FACTOR workflow.

### Creating Job Tokens

The first step to automating batch processing is to create the job tokens which hold the processing metadata and load them into the database. In order to track multiple concurrent reductions, these tokens are combined in sets. Once this set is given a name, a batch of tokens can be created for each pipeline step and uploaded into the PiCaS server. These tokens are set in the ‘todo’ state, indicating the processing has not yet started. The specific implementation of the ‘pre-FACTOR’ software is discussed in Appendix 3.A.

### Packing Scripts

While the PiCaS database can hold metadata which differs across jobs, pipelines or observations; the processing scripts need to be stored in a location where the worker machines have access to them. These scripts are archived and uploaded to storage and their location is added to the job token. After a worker machine locks a job token, it downloads and extracts the processing scripts, reads the metadata from the token and begins the processing (Fig.3.4).

The location of the scripts is stored in the job token. This allows different steps of the pipeline to use different sets of scripts. The benefits from this design is that as long as the worker node has access to the URI (Universal Resource Identifier) of the scripts, it can process the data, making data reduction portable over a variety of distributed computing environments, including the GRID. This capability will be used in the future to move processing to the location of the archive (Section 3.5.2). Implementation of this process is summarized in Appendix 3.A.

### Processing

Processing is launched with a launch script executed on a worker node. This script is responsible for locking a job token, taking it from the ‘todo’ state to the ‘locked’ state.

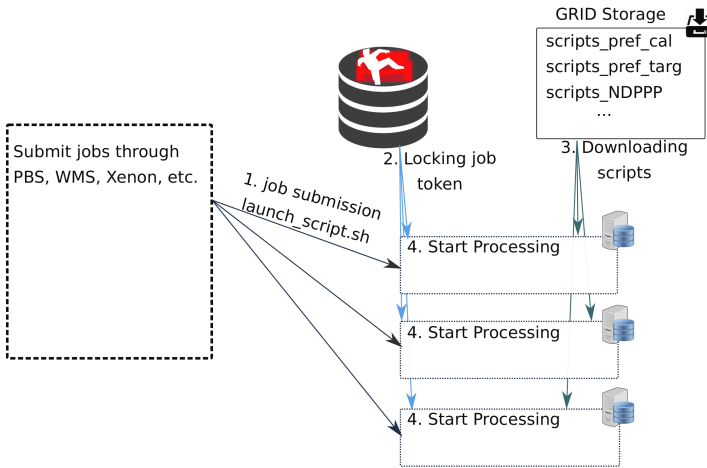


Figure 3.4: Starting processing on worker machines. Currently processing is done on SURF-sara Gina nodes, however the framework has been tested at the Leiden University cluster.

After the token is locked, the launch script downloads the script archive and launches the processing. For the LOFAR pre-FACTOR software, the scripts include the latest version of the pre-FACTOR repository. Additionally, there are helper scripts that set up the processing environment, download the data, post-process the output and upload the data to intermediate storage. The metadata stored in the PiCaS job token is fed into the setup and processing scripts.

Since the processing scripts and metadata are stored remotely, the same small launch script can load many different reduction pipelines simply by changing the group of tokens it should lock. This design choice makes it easy to create other script archives for the Direction Dependent pipeline or other LOFAR processing workflows.

### Intermediate Data Storage

Splitting the processing into multiple steps requires intermediate data to be stored at a location accessible to the worker machines. As the current processing is done at the Gina cluster, the intermediate results are stored in several dedicated storage pools hosted by SURFsara. The LRT processing scripts check for the availability of an initial data set or intermediate data product on launch and download it. This avoids unnecessary repetition of reduction steps and allows to restart a failed job.

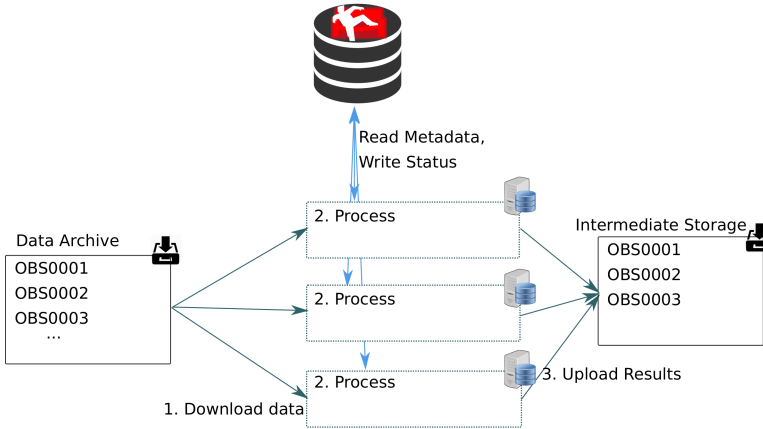


Figure 3.5: Processing of LOFAR data from the Long Term Archive with results stored at an intermediate storage location.

Since the data location is static, the paths of the data are hard coded into the scripts. The framework design, however, also allows a job to log the location of its output data into the CouchDB database. Jobs in subsequent steps can read the location of their input data from the tokens of the previous job. This will be implemented when the LOFAR data processing becomes distributed across multiple locations.

### 3.4.2 LOFAR Surveys Use Case

The LOFAR Two Meters Sky Survey requires the processing of more than 8 PB of data each year in order to keep up with the data produced by the telescope. As there are more than 3000 observations planned, processing them manually is untenable. Additionally, the large raw data sizes require the data be reduced in parallel before the Direction Dependent calibration step since the data will not fit in the memory of a single machine.

Re-purposing the LOFAR pre-FACTOR software to take advantage of the GRID computing resources by leveraging the automation provided by the LRT framework has allowed the processing of more than a hundred data sets in the time span of four months. Additionally, this was done in the time frame that an astronomer would normally produce less than 10 Direction Independent calibrated observations, hence a providing a necessary speed up by an order of magnitude.

### 3.4.3 Framework Capabilities

The LRT implementation is designed to be as platform-independent as possible. It allows for easy extensions enabling LOFAR reduction schemes other than the LoTSS reduction. Two examples of this are the updated LOFAR GRID spectroscopy and LOFAR GRID pre-processing pipelines (Oonk+ in prep). The GRID pre-processing pipeline runs flagging of bad data and averaging, however performs no calibration. The spectroscopy pipeline uses an independent set of scripts to perform calibration, bandwidth correction and imaging.

Thanks to the abstraction of the metadata and scripts storage, processing is also possible at other locations. The pre-FACTOR scripts require an installation of the LOFAR software stack[71]. These requirements are met by mounting a CernVM-FS[58][72] installation of the LOFAR stack. The CVMFS service provides a portable pre-compiled copy of the LOFAR software. With the CVMFS prerequisite satisfied and an active grid proxy, any computer can download data and process a job in any data reduction step.

### 3.4.4 Initial Results

Automatically launching jobs has made it possible to process more than 100 data sets between November 2016 and February 2017. Without a framework to automate and distribute the processing and a cluster at a Grid location, these data sets would need to be downloaded to an institute's cluster. Such standalone runs of the 'pre-FACTOR' scripts typically process one observation in two weeks, taking into account the data transfer time. At the 10MB/s connection (The sustained speed at Leiden University), the downloading would take over 5 years alone. Porting the LOFAR LoTSS data reduction to a Grid location using the LRT framework has resulted in a 15x increase in data throughput. Suggestions on further increasing the throughput are presented in Section 3.5.2.

## 3.5 Conclusion and Future Work

The goal of the LRT Framework is to create an automated software which can be used to port LOFAR processing to a massively distributed compute environment. The Direction Independent calibration of the LOFAR Two Meter Sky Survey was used as a demonstration of the capabilities of the LRT software.

Combining the ‘pre-FACTOR’ scripts with the LRT tools resulted in a 15x increase in throughput compared to previous data reduction strategies. Thanks to the automation provided, it is possible to process hundreds of observations, necessary for large astronomical survey projects such as the 3000+ observations of the LoTSS.

The scalability of this framework allows launching multiple data reductions concurrently and easily monitor their progress. The portability of the LRT framework makes it easy to move processing to archive locations, further increasing the throughput. Finally, as the framework is general, other LOFAR projects can increase throughput and automation by integrating their software.

### 3.5.1 Throughput Improvement

Using the LRT framework, more than 100 data sets have passed through the Direction Independent calibration. This is more than a 15 fold increase compared to the throughput at an institution with limited bandwidth to the data archive. From these data sets, 30 images have been produced since November 2016. The Direction Dependent pipeline, responsible for the imaging of the DI calibrated data, is not yet automated using the LRT framework. The topic of this automation will be handled in future work. Future improvements (Section 3.5.2) are expected to increase the throughput to one image per day.

Currently, data reduction is launched manually. There are upcoming plans to create a trigger launched by the Observatory at the end of a successful observation. Using this trigger, the processing can be integrated with the data acquisition and launched automatically at the end of the observation. Doing so enables producing an image less than a week after the observation has completed without requiring human interaction.

### 3.5.2 Future Work

Most of the LoTSS data is not stored at the SURFsara Grid location. Because of the high data sizes, the 1 Gbps transfer between these remote sites and SURFsara is insufficient to process the 3000+ data sets. The proposed solution is to launch the initial reduction steps (Calibrator 1 and Target 1 in Fig.3.2) on compute clusters at the archive locations. Running the Target 1 step at these sites will reduce the data size from 16TB to 0.5TB. The resulting intermediate data can easily be transferred

over a 1 Gbps connection in 1 hour compared with the 35 hours for the original data.

While the LRT framework successfully automated the Direction Independent calibration pipeline, it still needs to implement the Direction Dependent processing scripts. The DI data can be easily split into pieces and processed concurrently, however current Direction Dependent pipelines cannot split the data easily. This means they will only benefit from the automation aspect of the framework. Because of this limitation of the algorithms, (Direction Dependent) processing of each data set needs to be done on a single machine. This part of the processing currently takes more than four days per data set per node. To process all 3000 observations within five years, the DD reduction will need at least eight dedicated machines continuously processing data. Nevertheless, the input data is only 200-500GB making it easier to store and transfer to institutes not part of the European Grid Infrastructure.

The Xenon framework[73] allows launching jobs at multiple clusters from a single location. Thanks to the portability of the LRT software, it will be possible to integrate the LOFAR reduction with Xenon. This will automate the launching of Direction Dependent processing at multiple institutions. Automatically launching these jobs will make efficient use of the computer resources at SURFsara and other sites. Using this strategy, the LoTSS project data can be processed within the anticipated time.

Finally, as the reduction is automated, it can be started right after the telescope finishes the observation. Launching jobs immediately after an observation will minimize the time staging the data from tape to disk. Currently, the staging process can take up to a week. Triggering the processing immediately after observation, an image will be produced less than a week after the data is acquired.

Minimizing the latency between observation and science quality images will benefit the LOFAR community immensely by allowing radio astronomers to focus on their specific science case. An all-sky survey at the 150 MHz range will create a multitude of targets for follow-up with optical telescopes and result in many discoveries in the field of Radio Astronomy. A full list of science results expected from the LoTSS project can be found in[25]. Efficient high-throughput processing of LOFAR data will empower the above science cases opening the way to exciting discoveries.

## 3.A Execution of LOFAR Reduction Tools

The LRT framework handles staging of LOFAR data, packaging and uploading worker node scripts (named sandboxes), creating PiCaS job tokens and launching pilot jobs on the SURFsara Gina cluster. The framework is modular, allowing a user to execute any of the previous steps manually, or alternatively launch an automated reduction. It can be downloaded from the Github page ([github.com/apmechev/GRID\\_LRT](https://github.com/apmechev/GRID_LRT)) and installed with

```
python setup.py build && python setup.py install .
```

Documentation on the usage of the tools can be found at the Github page.



**Original Abstract:** Modern astronomical data processing requires complex software pipelines to process ever growing data sets. For radio astronomy, these pipelines have become so large that they need to be distributed across a computational cluster. This makes it difficult to monitor the performance of each pipeline step. To gain insight into the performance of each step, a performance monitoring utility needs to be integrated with the pipeline execution. In this work, we have developed such a utility and integrated it with the calibration pipeline of the Low Frequency Array, LOFAR, a leading radio telescope. We tested the tool by running the pipeline on several different compute platforms and collected the performance data. Based on this data, we make well informed recommendations on future hardware and software upgrades. The aim of these upgrades is to accelerate the slowest processing steps for this LOFAR pipeline. The *pipeline\_collector* suite is open source and will be incorporated in future LOFAR pipelines to create a performance database for all LOFAR processing.

This chapter is based on the work of **A.P. Mechev**, A. Plaat, et al. “Pipeline Collector: Gathering performance data for distributed astronomical pipelines”. In: *Astronomy and Computing* 24 (2018), pp. 117–128. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2018.06.005>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133718300490>

## 4.1 Introduction

Astronomical data often requires significant processing before it is considered ready for scientific analysis. This processing is done increasingly by complex and autonomous software pipelines, often consisting of numerous processing steps, which are run without user interaction. It is necessary to collect performance statistics for each pipeline step. Doing so will enable scientists to discover and address software and hardware inefficiencies and produce scientific data at a higher rate. To identify these inefficiencies, we have extended the performance monitoring package *tcollector*<sup>1</sup>[74]. The resulting suite, *pipeline\_collector*, makes it possible to use *tcollector* to record data for complex pipelines. We have used a leading radio telescope as the test case for the *pipeline\_collector* suite. The discoveries made with our software will help remove bottlenecks and suggest hardware requirements for current and future processing clusters. We summarize our findings in Table 4.1 in Section 4.3.

Over the past two decades, processing data in radio astronomy has increasingly moved from personal machines to large compute clusters. Over this time, radio telescopes have undergone upgrades in the form of wide-band receivers and upgraded correlators [75, 76]. In addition, several aperture synthesis arrays such as the Low Frequency Array [LOFAR, 77], Murchison Widefield Array [MWA 78, 79] and MeerKAT [80] have begun observing the radio sky, leading to an increase of data rates by up to 3 orders of magnitude [81, 82].

As the data acquisition rate has increased, data size has entered the Petabyte regime, and processing requirements increased to millions of CPU-hours. In order for processing to match the acquisition rate, the data is increasingly processed at large clusters with high-bandwidth connections to the data. An important case where data processing is done at a high throughput (HTC) cluster is the LOFAR radio telescope.

The LOFAR telescope is a European low frequency aperture synthesis radio telescope centered in the Netherlands with stations stretching across Europe. This aperture synthesis telescope requires significant data processing before producing scientific images [17, 18, 83, 84]. In this work, we will use our performance monitoring utility, *pipeline\_collector*<sup>2</sup>, to study the first half of the LOFAR processing, the Direction Independent (hereafter DI) pipeline.

<sup>1</sup><https://github.com/OpenTSDB/tcollector>

<sup>2</sup>[https://gitlab.com/apmechev/pipeline\\_collector.git](https://gitlab.com/apmechev/pipeline_collector.git)

One major project for the LOFAR telescope is the Surveys Key Science Project (SKSP) [25]. This project consists of more than 3000 observations of 8 hours each, 600 of which have been observed. These observations need to be processed by a DI pipeline, the results of which are calibrated by a Direction Dependent (DD) pipeline. The DI pipeline is implemented in the software package *prefactor*<sup>3</sup>. The *prefactor* pipeline is itself split into four stages and implemented at the SURF-sara Grid location at the Amsterdam e-Science centre [35, 85]. The automation and simple parallelization has decreased the run time per data set from several days to six hours, making it comparable to the observation rate. To better understand and optimise the performance of the *prefactor* pipeline, we require detailed performance information for all steps of the processing software. We have developed a utility to gather this information for data processing pipelines running on distributed compute systems.

In this work, we will use the *pipeline\_collector* utility to study the LOFAR *prefactor* pipeline and suggest optimisation based on our results. To test the software on a diverse set of hardware, we will set up the monitoring package on four different computers and collect data on the pipeline's performance. Using this data, we discuss several aspects of the LOFAR software which we present in Table 4.1. Finally, we discuss the broader context of these optimisations in relation to the LOFAR SKSP project and touch on the integration of *pipeline\_collector* with the second half of the data processing pipeline, the DD calibration and imaging.

#### 4.1.1 Related Work

Scientific fields that need to process large data sets employ some type of data processing pipelines. Such pipelines include e.g. solar imaging [86], neuroscience imaging [87] and infrared astronomy [88]. While these pipelines often log the start and finishing times of each step (using tools such as *pegasus-kickstart* [89]), they do not collect detailed time series performance data throughout the run.

At a typical compute cluster the performance of every node in a distributed systems is monitored using utilities, such as Ganglia [90]. These tools only monitor the global system performance. If one is interested in specific processes, then the Linux *procfs* [91] is used. The *procfs* system can be used to analyze the performance of individual pipeline steps. Likewise, the Performance API [PAPI, 92] is a tool which collects detailed low level information on the CPU usage per process. Collecting detailed statistics at the process level is required to understand and

<sup>3</sup>available at <https://github.com/lofar-astron/prefactor>

| Result #  | Description                                                                                                                          |
|-----------|--------------------------------------------------------------------------------------------------------------------------------------|
| <b>R1</b> | Native compilation of the software performs comparably to pre-compiled binaries on two test machines.                                |
| <b>R2</b> | The processing steps do not appear to accelerate significantly on a faster processor or with larger cache size.                      |
| <b>R3</b> | Both calibration steps ( <i>calib_cal</i> and <i>gsmcal_solve</i> ) show linear correlation between speedup and memory bandwidth.    |
| <b>R4</b> | Disk read/write speed does not affect the completion time of the slowest steps.                                                      |
| <b>R5</b> | Both calibration steps do not use large amounts of RAM despite processing data on the order of Gigabytes.                            |
| <b>R6</b> | The <i>calib_cal</i> step can suffer up to 20% of Level 1 Instruction Cache misses, while <i>gsmcal</i> only has 5% of these misses. |
| <b>R7</b> | Both calibration steps are impacted by Level 2 Cache eviction at comparable rates.                                                   |
| <b>R8</b> | The <i>calib_cal</i> step stalls on resources 70% of cycles while the <i>gsmcal</i> step only 30% of them.                           |
| <b>R9</b> | The <i>calib_cal</i> uses the CPU at full efficiency for only 10 % of the CPU cycles.                                                |

Table 4.1: A table of all the results presented in Section 4.3.

optimise the performance of the LOFAR pipeline and we will integrate PAPI into *pipeline\_collector* in the future. Finally, DTrace[93] is a Sun Microsystems tool which makes it possible to write profiling scripts that access data from the kernel and can be used to monitor process or system performance at run time with minimal overhead. As DTrace was not installed on either of the processing clusters, we have not used it to monitor the pipeline's performance.

The Linux procs system and PAPI record data which is already made available by the Linux kernel. This option incurs insignificant overhead as it uses data the kernel and processor already log. Likewise PAPI reads performance counters that the CPU automatically increments during processing. These profiling utilities can run concurrently with the scientific payload without using more than 1-2% of system resources. Their low overhead is why we choose to use them to collect performance data.

Other tools for performance analysis such as Valgrind [94] collect very detailed performance information. This comes at the expense of execution time: running with Valgrind, the processing time slows by up to two orders of magnitude. As such, we do not use Valgrind along the LOFAR software.

## 4.2 Measuring LOFAR Pipeline performance with *pipeline\_collector*

We developed the package *pipeline\_collector* as an extension of the performance collection package *tcollector*. *pipeline\_collector* makes it possible to collect performance data for complex multi-step pipelines. Additionally, it makes it easy to record performance data from other utilities. A performance monitoring utility that we plan to integrate in the future are the PAPI tools described in section 4.1.1. The resulting performance data was recorded in a database and analyzed. For our tests, we used the LOFAR *prefactor* pipeline, however with minor modifications, any multi-step pipeline can be profiled.

*tcollector* is a software package that automatically launches 'collector' scripts. These scripts sample the specific system resource and send the data to the main *tcollector* process. This process then sends the data to the dedicated time series database. We created custom scripts to monitor processes launched by the *prefactor* pipeline (4.A.1).

In this work, we use a sample LOFAR SKSP data set as a test case. A particular focus was to understand the effect of hardware on the bottlenecks of the LOFAR data reduction. To gain insight into the effect of hardware on *prefactor* performance, the data was processed on four different hardware configurations (Table 4.2). As typical upgrade cycle for cluster hardware is five years, our results will be used to select optimal hardware for future clusters tasked with LOFAR processing.

### 4.2.1 *Prefactor* Pipeline

The LOFAR *prefactor* pipeline [18] is a software pipeline that performs direction independent calibration using the LOFAR software. The LOFAR software stack is a software package containing commonly used processing software used by LOFAR pipelines [44, 95]. These tools are built and maintained by ASTRON<sup>4</sup>.

The *prefactor* pipeline performs a sequence of four stages, namely the calibrator and target calibration. The first half of *prefactor* processes data from a calibration source and the second half processes a science target. Altogether, this processing takes six hours on a high-throughput cluster. The final result is a data-set ready for

<sup>4</sup>ASTRON: Netherlands Institute for Radio Astronomy, url<https://www.astron.nl/>

creating images of the sky at radio wavelengths. Figure 4.1 shows a graphical view of the prefactor pipeline’s Calibrator and Target stages.

The Calibrator stage consists of the *ndppp\_prep\_cal* and the *calib\_cal* step. The former flags radio interference and averages the data, and the latter performs gain calibration on a bright calibration source. It is followed by the *fitclock* step which fits a clock-TEC model to the calibration solutions [18].

The Target stage consists of a *ndppp\_prep\_targ* step, *predict\_ateam*, *gsmcal\_solve* and *gsmcal\_apply* steps. The first two of these steps flag and average the target data and calculate contamination by bright off-axis radio sources. The *gsmcal\_solve* step determines phase solutions for each antenna using a model of the target and the results of the *ndppp\_prep\_targ* step. Finally, the *gsmcal\_apply* step applies these solutions to the target data. Figure 4.2 shows the percentage of time spent by these steps for the four *prefactor* stages.

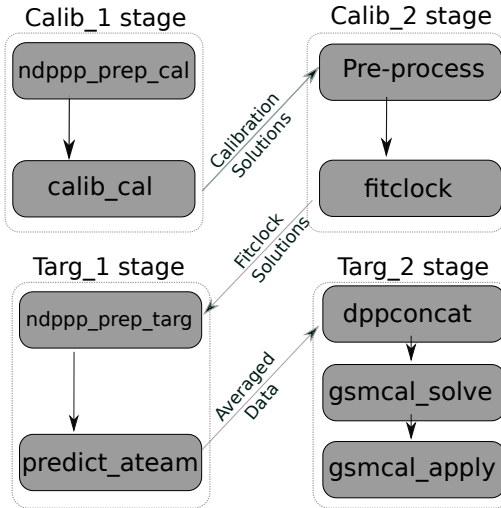


Figure 4.1: The four processing stages that make up the prefactor pipeline. The Calibrator stages (top) process a known bright calibrator to obtain the gain for the LOFAR antennas. The Target stages (bottom) process the scientific observation to remove Direction Independent effects. The *pref\_cal1* and *pref\_targ1* stages are massively parallelized across nodes without the need for an interconnect. The *pref\_cal2* step runs only on one node, while *pref\_targ2* is parallelized on 25 nodes. As the LOFAR software does not use MPI, we can run each processing job in isolation.

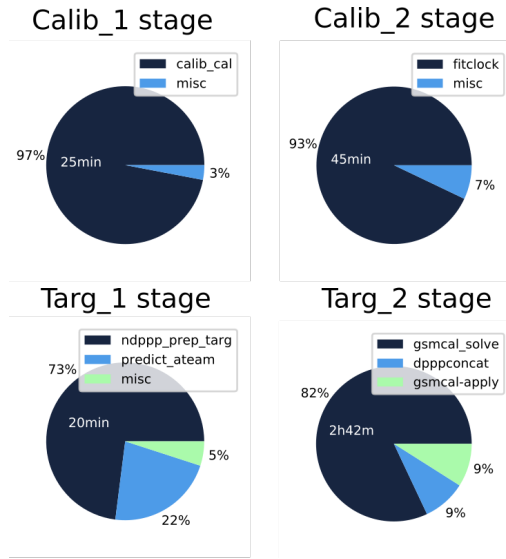


Figure 4.2: Portion of processing time taken by each step for the four prefactor stages, as reported by the Prefactor software. For each stage, the majority of processing time is spent during one or two steps. This is due to the fact that each prefactor stage also has intermediate book-keeping steps explicitly included in the pipeline. For each pipeline stage, the mean processing time for the longest-running steps at SURFsara is also indicated. It should be noted that while faster, *pref\_cal1* runs ten times as many jobs as *pref\_targ2*.

#### 4.2.2 Performance suite

Cluster performance is frequently monitored using utilities such as Ganglia [90, discussed in Section 4.1.1]. These tools cannot access individual processes and thus cannot collect data on a per-process basis. To collect such data, each process launched by the active pipeline step needs to be profiled individually. Our utility is designed to gather such performance data.

Our monitoring package, *pipeline\_collector* adds custom performance collectors (4.A.1) to the performance collection framework *tcollctor*. We use these collectors to monitor individual pipeline steps as defined by the user<sup>5</sup>. The tools attach to processes launched by the pipeline and record performance data at a one second interval. This sampling frequency is at high enough resolution to detect trends and anomalies in hardware utilization, and still result in a reasonable database size. The performance data is uploaded to a remote time series database, OpenTSDB [96]. Details on the data collection can be found in 4.A.

<sup>5</sup>[https://gitlab.com/apmechev/pipeline\\_collector.git](https://gitlab.com/apmechev/pipeline_collector.git)

| Location | CPU Speed (MHz) | CPU Model | Micro-architecture | Cache Size | RAM Speed <sup>7</sup> | Disk Speed <sup>8</sup> |
|----------|-----------------|-----------|--------------------|------------|------------------------|-------------------------|
| Leiden   | 2200            | E5-4620   | Sandy Bridge       | 16 MB      | 1.4 GB/s               | 99.7 MB/s               |
| SURFsara | 2500            | E5-2680   | Sandy Bridge       | 30 MB      | 2.5 GB/s               | 65.4 MB/s               |
| Hatfield | 2900            | E5-2660   | Sandy Bridge       | 20 MB      | 2.4 GB/s               | 155 MB/s                |
| Laptop   | 3300            | E3-1505M  | Skylake            | 8 MB       | 4.7 GB/s               | 822 MB/s                |

Table 4.2: CPU, Cache, RAM and Storage specifications of the four test machines. The tested machines span a factor of 1.5x in CPU speed, 4x in cache and RAM Speed and 10x in Disk speed.

## Performance API

The time-series database is also used to collect low-level CPU information for each process. This information is collected by the PAPI interface (discussed in Section 4.1.1). This was done through the `papiex` utility<sup>6</sup> [97]. This utility records the CPU’s internal performance counters. A CPU’s performance counters record information on how efficiently the software uses the CPU’s resources. The results from this test are detailed in Section 4.4.

## 4

### 4.2.3 Test Hardware

In order to study the effect of different hardware configurations on the performance of LOFAR processing, the *prefactor* pipeline was run on four different sets of hardware. The four machines tasked with processing LOFAR data comprised nodes at three computational clusters and a personal computer. The tests were run while the systems were idle to make sure there is no interference of other software with the LOFAR processing. Table 4.2 details the specifics of the four test machines.

## 4.3 LOFAR Prefactor Test Case

With the test set described in Section 4.2, we aim to understand processing bottlenecks in the *prefactor* pipeline and make informed decisions on future hardware and software upgrades. To do so, we processed a sample observation at institutes that typically process LOFAR data.

From the data collected by processing the sample observation, we determined the slowest pipeline steps. These steps were the *calib\_cal* and *gsmcal\_solve*, seen in Figure 4.2. The *calib\_cal* step is implemented by the software `bbs-reducer`

<sup>6</sup>Available at <https://bitbucket.org/minimalmetrics/papiex-oss>

[95, 98] and the *gsmcal\_solve* step is implemented by NDPPP [95, 99]. Both *bbs-reducer* and NDPPP are part of the LOFAR software suite.

We collected performance statistics using the *pipeline\_collector* suite as discussed in Section 4.2. The run time of the slowest *prefactor* steps on the four machines is shown in figure 4.3. The results discovered using *pipeline\_collector* are listed in Table 4.1 and discussed in Section 4.3.2. Using the PAPI interface (discussed in Section 4.2.2) CPU performance data was collected. The results from this test are detailed in Section 4.4.

We will present a number of insights into the performance of the LOFAR software collected by the profiling suite. The results are presented in Table 4.1 and are grouped in three main areas. The effect of compilation on the run time was result **R1**. The set of results **R2**, **R3**, **R4** and **R5** were obtained using the *pipeline\_collector* package. Results **R6**, **R7**, **R8** and **R9** were collected with the PAPI package, which will be integrated into *pipeline\_collector* in the future.

### 4.3.1 Pre-compiled vs native compilation

The performance trade-off between pre-compiled and native compilation was studied first. The majority of the processing for the LOFAR SKSP Project [25] is done at the SURFsara GINA cluster in Amsterdam. This location is part of the European Grid Initiative (EGI)[85]. At this location, the software is deployed by compiling on a virtual machine and mounting it on all worker nodes through the CernVM FileSystem (CVMFS) service [100]. The CVMFS server allows any client to mount a fully compiled LOFAR installation, making it easy to distribute and version control the software within and outside of SURFsara. An alternative is to locally compile the LOFAR packages on each cluster. The performance of the natively compiled<sup>9</sup> vs CVMFS installations was compared on the laptop test machine using *pipeline\_collector*. In order to validate this result, the two compilations of the same software were also tested at the Data Science Lab at the Leiden Institute of Advanced Computer Science (LIACS)<sup>10</sup>.

<sup>7</sup>benchmarked using *dd*

<sup>8</sup>sequential disk read, benchmarked using *fio* - flexible I/O tester:

```
fio -randrepeat=1 -ioengine=libaio -direct=1 -gtod_reduce=1 -name=test -filename=test
-bs=4k -iodepth=256 -size=4G -readwrite=read -ramp_time=4
```

<sup>9</sup>The software was compiled using `-march=native` and `-O3` compilation flags. On the laptop, gcc resolves `-march=native` as `broadwell`. The CVMFS installation resolves `-march=native` as `core-avx-i`.

<sup>10</sup><https://www.universiteitleiden.nl/en/science/computer-science/about-us/our-facilities>

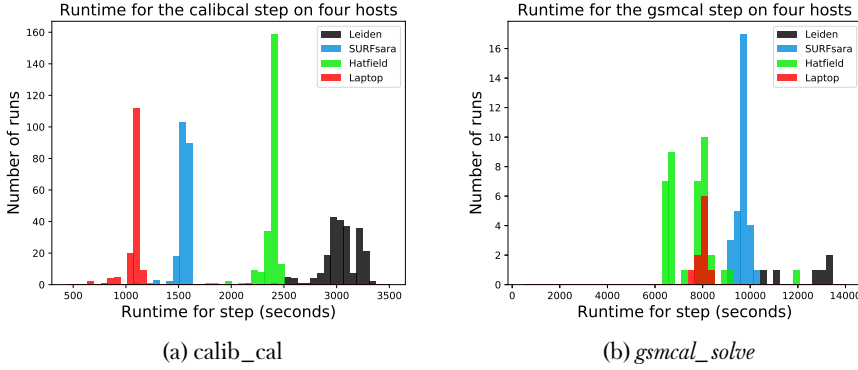


Figure 4.3: Job completion times for *calib\_cal* and *gsmcal\_solve* steps tested on four hardware setups. The *calib\_cal* step ran 244 times. The *gsmcal\_solve* ran 24 times as the data is concatenated from 244x1 to 24x10 sets. The step with the longest latency is *gsmcal\_solve* while *calib\_cal* consumes a comparable number of core-hours over 244 jobs.

4

An interesting discovery is that the LOFAR software did not process data faster when compiled natively. This is despite the fact that the local install was compiled with advanced processor instructions available on the host machine. Figure 4.4 shows a histogram of its processing time with the two different compilation options for the *calib\_cal* software running on the sample data set. The same test was done for the software performing the gain calibration (*gsmcal\_solve*), seen in Figure 4.5. The result of this experiment is shown in Figures 4.4a and 4.5a. The software compiled at SURFsara showed a minor improvement for the *calib\_cal* step on the laptop machine, however this improvement is not seen on the computational cluster node.

Overall, the software for both steps shows no significant improvement when compiled natively. This is result R1 in Table 4.1. The second run at LIACS also confirms this result for both steps (Figures 4.4b and 4.5b). This result suggests that the slowest *prefactor* steps are not optimised for modern processors.

### 4.3.2 Prefactor Run time and Hardware Parameters

Next, we studied the dependence of run time on different hardware parameters. With software that collects per-step performance statistics for the LOFAR pipeline, the dependence of the pipeline processing on hardware performance can be easily profiled and studied. Using *pipeline\_collector* we determined the pipeline's slowest steps with respect to different hardware parameters.

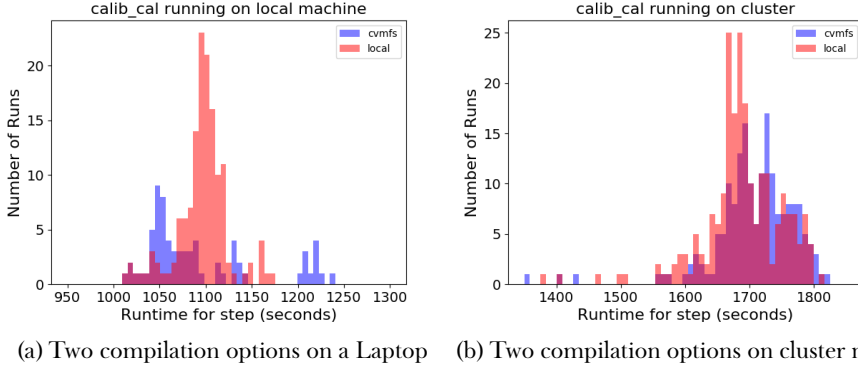


Figure 4.4: Difference in processing time for *calib\_cal* when compiled remotely and natively. *calib\_cal* was run 244 times with the native software and 40 times with the CVMFS compilation. Two tests were done: one on the personal laptop (4.4a) and one on a cluster node at the LIACS Data Science Lab (4.4b). The test on a cluster node shows no significant difference in run time between compilation options. The laptop test suggests that the remotely compiled software may run 5% faster than the local compilation.

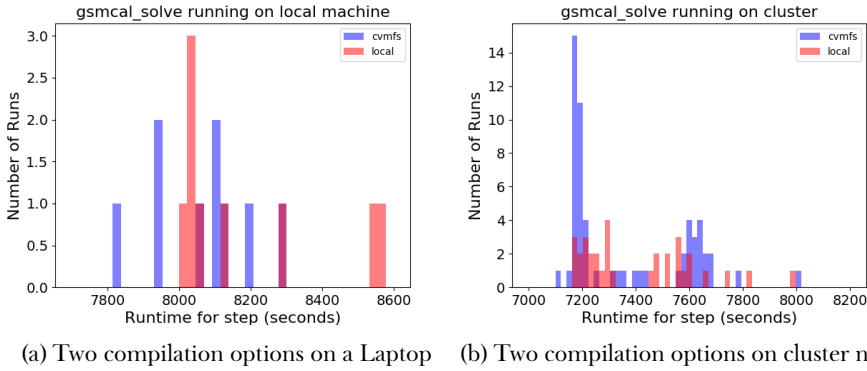


Figure 4.5: Difference in processing time for *gsmcal\_solve* when compiled remotely and natively. *gsmcal\_solve* was run 50 times with the native software and 120 times with the CVMFS compilation. Two tests were done: one on the personal laptop (4.5a) and one on a cluster node at the LIACS Data Science Lab (4.5b). Just like with the *calib\_cal* step, the *gsmcal\_solve* step also doesn't accelerate significantly when natively compiled.

The system parameters studied here are the CPU speed, memory throughput, cache size and disk speed. Modern computers can have a complex memory hierarchy, as demonstrated in Figure 4.6 [101]. This is due to the cost trade-off between memory size and memory speed. Because of this trade-off, the full data set is stored on disk, while the working set is placed in RAM. This is the data that the processor needs to access at the current time [102]. The most frequently accessed parts of the data are stored in the CPU cache, which evicts the oldest data when full [103].

The CPU processing speed is faster than the RAM latency, so a hierarchy of caches exist. Caches store small subsets of the working set and have a fast connection to the processor. The fastest data link is between the CPU and the L1 Cache, with the link to RAM being slower and the disk read speed slower still. The limited memory capacity of the different levels of the memory hierarchy as well as the throughput between them will lead to performance bottlenecks. These bottlenecks will lead to the processor waiting on memory. Such stalls lead to longer processing times.

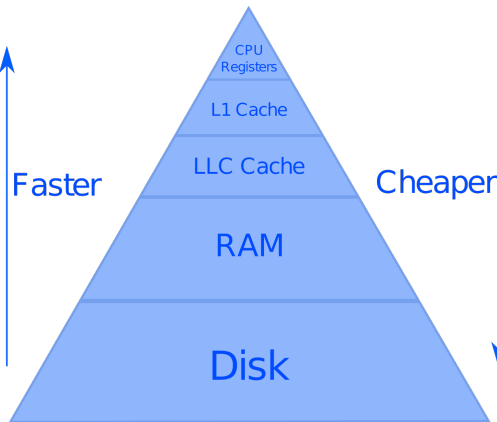


Figure 4.6: A model of the memory hierarchy, as described in [104].

## CPU

The CPU speed is usually the primary factor determining how fast computations can be made. In general, a faster CPU will result in faster data processing.

However, Fig. 4.7a shows that the run time of the calibration of the calibrator does not strongly depend on the CPU frequency. While the test nodes at SURFsara and Leiden run at the same CPU frequency, running on a cluster node at

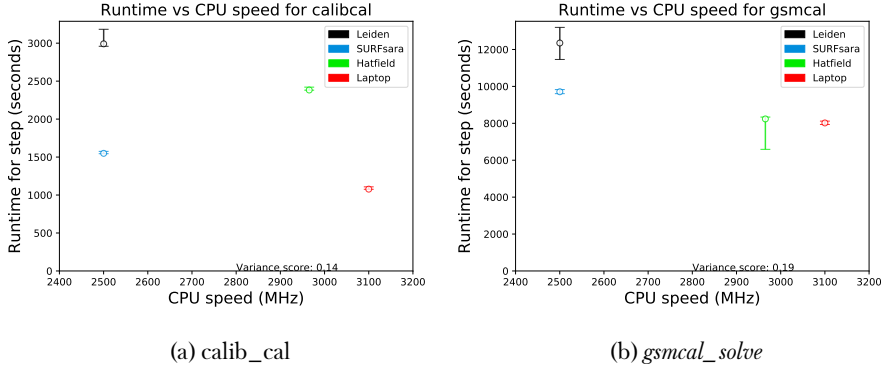


Figure 4.7: Performance of the bottleneck steps compared with the CPU speeds of the four test machines. The values are the mean of 244 runs (Standard prefactor run) and the error bars show the 1-sigma of the distribution of the run time.

SURFsara takes half the time as on a node at Leiden. Even more surprisingly, the *gsmcal\_solve* step does not benefit significantly from a faster CPU, despite being the most computationally heavy *prefactor* step (**R2**). This step does the gain calibration on the target field using the StEFCal algorithm [105]. Figure 4.7b shows only a slight improvement over faster CPU clock speeds for both steps. The correlation between completion time and CPU speed is similar for both steps.

## Cache

The CPU has a hierarchy of caches consisting of Level 1, Level 2 Cache and LLC Cache. For the four processors tested, the Level 1 and 2 caches were all the same size, thus the only difference is the Last Level Cache (LLC or just Cache in Figure 4.6). This cache stores data needed by the CPU, so the larger it is, the less the processor needs to wait for RAM to return data.

In general, numerical codes benefit from larger cache sizes [104, 106]. Interestingly, figure 4.8b suggests that the *gsmcal\_solve* step does not exclusively depend on larger cache **R3** (Table 4.1). On the machines with a larger cache, the *gsmcal\_solve* step completed processing as quickly as on the machines with smaller cache, even down to 8MB.

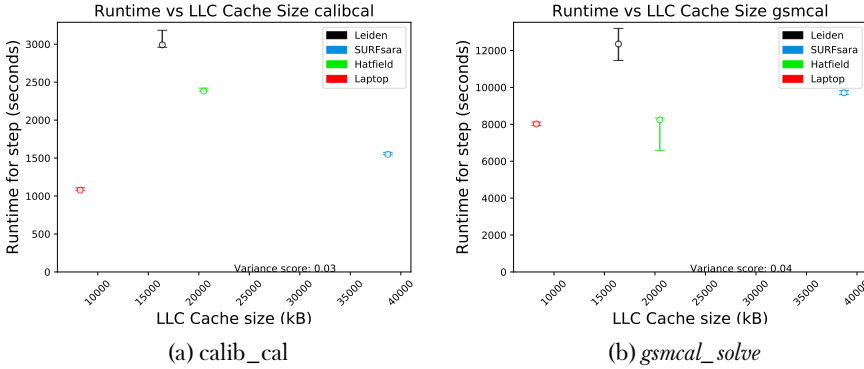


Figure 4.8: Performance of the two bottleneck steps with respect to Last Level Cache size. The *gsmcal\_solve* step shows no trend between cache size and completion time. The *calib\_cal* step runs the fastest on the machine with the smallest cache.

## RAM Bandwidth

4

If the entire data set does not fit into cache, the software needs to transfer data from RAM to the CPU. In these cases, *prefactor* benefits from a fast bandwidth between the cache and RAM. For this study, the RAM throughput was benchmarked<sup>11</sup>. This command copies dummy data into system memory. As this utility exists on all Unix systems, this is a standardized benchmark of the RAM performance.

Figure 4.9a showed that higher bandwidth is correlated with a faster completion time for the *calib\_cal* and *gsmcal\_solve* steps (**R4**). The result is to be expected as the working set of these steps is 200MB and 1.0GB respectively, and cannot fit into cache readily, however it is loaded into RAM within the first 5 seconds of the run (Figure 4.10), and is streamed from memory throughout the run.

## Disk Read speeds

The slowest link in the memory hierarchy is the disk read speed. For the *calib\_cal* step, the entire data is loaded into memory during the first few seconds of the run, after which the disk only becomes important when the results need to be written out. The *gsmcal\_solve* step streams data from the disk to memory throughout the entire run. The plot of disk read speeds (Fig. 4.11b) also shows that a faster disk does not speed up the slowest step **R5**. To verify that disk throughput was not the limiting

<sup>11</sup>Using the command `$> dd if=/dev/zero of=/dev/shm/test bs=1M count=2048`

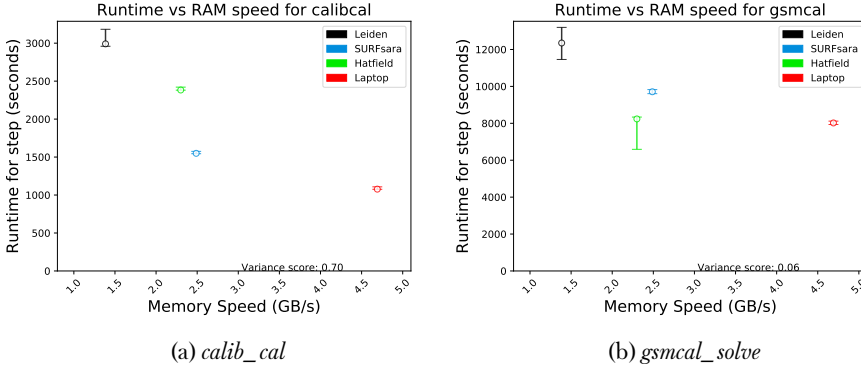


Figure 4.9: Performance of the two bottleneck steps and RAM bandwidth in GB/s. Both the *calib\_cal* and *gsmcal\_solve* steps show a trend of faster processing times on machines with higher RAM bandwidth. Both steps show a trend of decreasing processing time with increasing RAM throughput.

factor, the entire data set (25 GB) was moved to main memory (using `/dev/shm`). The resulting run time for both bottleneck steps did not change.

The calibration steps both stored less than 200MB of data in memory throughout their run. Figure 4.10 shows the time-series of the total memory used by these steps. The *calib\_cal* step uses only 200MB of memory and *gsmcal\_solve* only 35MB. While the *gsmcal\_solve* step works on a 1GB data set, it streams the data in memory and thus does not require 1GB of RAM. Alternatively, the *calib\_cal* step loads the entire (200MB) data set into memory for the entire duration of the run. The RAM usage time-series in Figure 4.10 show that the RAM is filled for the first 5 seconds of the run, further confirming that the processing is effectively independent from disk speed.

## 4.4 CPU Utilization Tests with PAPI

To gain more fine grained data on the CPU utilization, the *calib\_cal* and *gsmcal\_solve* steps were tested with the PAPI package. We ran this package as a test, to determine whether collecting PAPI data is helpful in understanding pipeline performance. PAPI can record data such as cache performance, branch prediction rate, fraction of memory/branch instructions and others. This data is complementary to the `procs` information, which is collected by the Linux kernel. As the collected data was useful in understanding the *prefactor* pipeline, we will include PAPI in the *pipeline\_collector*

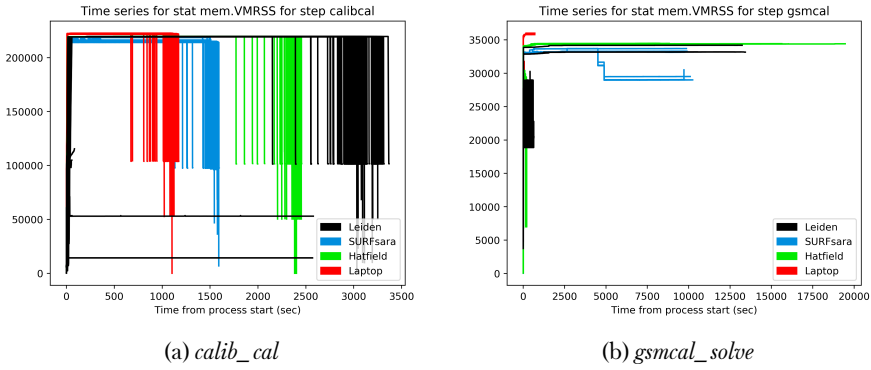


Figure 4.10: Time series of the Virtual Memory Resident Set Size. This is the amount of data stored in RAM (in Kb) during the *calib\_cal* and *gsmcal\_solve* steps. Both steps show the same amount of memory use on all test machines. Additionally, after a brief loading of data, the memory usage remains constant until processing is finished.

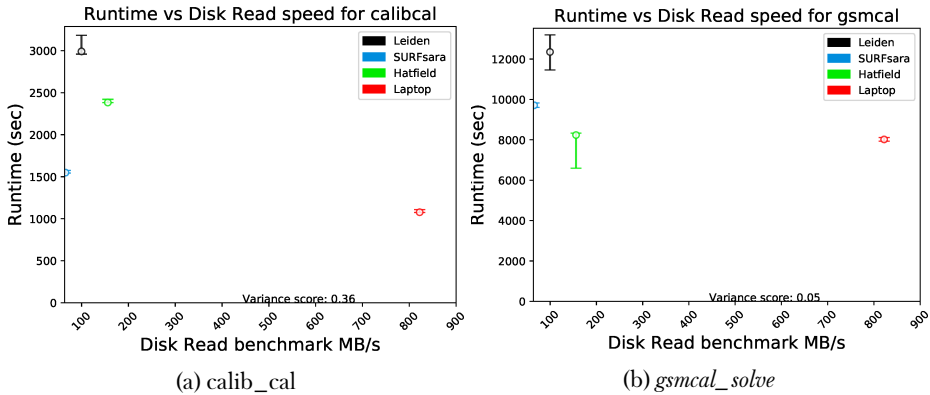


Figure 4.11: Performance of the two bottleneck steps and Disk bandwidth in MB/s. There is no correlation between the Disk read speed and the Run time of the steps.

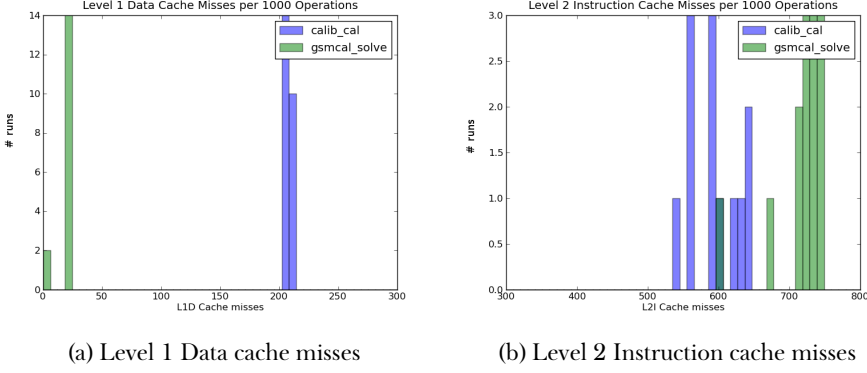


Figure 4.12: Cache miss rates for *calib\_cal* and *gsmcal\_solve*, executed on the SURFsara gina cluster. The cache is split into instruction and data caches. The figures above show the difference in the number of cache misses for both instruction cache and data cache for the slowest *prefactor* steps. *calib\_cal* suffers significantly more Data cache misses than *gsmcal\_solve* while the two steps undergo similar instruction Cache misses.

suite in the future. In the following sections we will discuss the results obtained for the *calib\_cal* and *gsmcal\_solve* steps.

#### 4.4.1 Level 1 Data Misses

The Level 1 Cache is split into cache for instructions and data. For all our test hardware the L1 Data cache is 32 Kb, and has a direct link to the processor's computational units [107]. The processor collects information logging how many times data requested by the CPU is not located into the L1 Data cache. This counter is called the Level 1 Data Cache Miss rate. To resolve this type of cache miss, the data needs to be fetched from L2 Cache. When this happens, the processor has to wait for the requested data. **R7:** The recorded L1 data misses in Figure 4.12a, show that the software performing the *calib\_cal* step misses 20% of its L1 data cache requests, while the software implementing the *gsmcal\_solve* step misses less than 5% of L1 Cache requests. These cache misses often happens in multi-threaded applications where there are instructions shared by multiple threads on the same cache line [108].

#### 4.4.2 Level 2 Instruction Misses

Unlike the Level 1 cache, Level 2 cache stores data and instructions in the same location. When the cache is full, it evicts the last used element in order to make

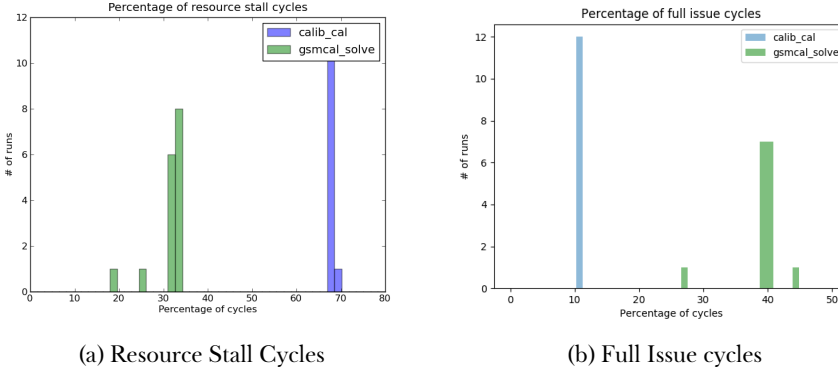


Figure 4.13: Resource stall cycles and Full Instruction Issue cycles. The two steps were executed on the SURFsara gina cluster.

space for newly requested data. PAPI also counts these eviction events. Figure 4.12b shows that for both steps, between 50 and 70% of L2 requests for an instruction do not match the contents of L2 Cache. This is significantly more than the applications benchmarked in [109, Table 2]. Because both steps process data of considerable size, the large amount of data required can evict instructions from the L2 cache (insight number **R7** in table 4.1).

#### 4.4.3 Resource Stalls

Modern processors have multiple computational pipelines on chip, in order to process data in parallel [110]. There are times when the processor's internal pipeline needs to wait for other instructions to finish. When this happens, it flags that it has 'stalled on a resource'. These resource stall cycles are also recorded by PAPI and represented as a percentage of total cycles. From figure 4.13a, it can be seen that *calib\_cal* stalls on 70% of the processor cycles, while *gsmcal\_solve* only on 33% of cycles (**R8**).

The Full Issue Cycles counter indicates the percentage of processor cycles, in which the theoretical maximum number of instructions are executed. During these cycles, the software uses the CPU optimally. The full issue cycles counter (Fig. 4.13b) also shows the difference in efficiency between the *calib\_cal* and *gsmcal\_solve* step (**R9**), with the former only working at peak efficiency for 10% of the processor cycles.

The plots in Figures 4.13a and 4.13b indicate that the *calib\_cal* step does

not use the internal CPU pipelines efficiently leading to waiting on resources and sub-optimal use of the CPU's Computational Units.

## 4.5 Discussions and Recommendations

With an increase of data acquisition rates and data complexity in radio astronomy, it is becoming important to thoroughly understand and optimise the performance of processing pipelines. Using *pipeline\_collector*, data can be collected for each pipeline step without altering the processing software. We store this data in a time-series database. The collected data can be studied to help researchers understand the pipeline performance for different processing parameters, data sets, and on different hardware. The *pipeline\_collector* suite is easy to deploy for mature pipelines and has minimal impact on pipeline performance. Typical CPU usage is  $<0.2\%$  with a memory footprint of  $\sim 1\text{-}10$  MB.

Creating a performance model with the collected data will allow us to optimise future clusters for LOFAR data processing. Doing so is necessary given the current data throughput, number of observations and time-line of the SKSP project. Similar issues will be encountered with upcoming radio telescopes [111].

To showcase the power of the *pipeline\_collector* suite, the LOFAR *prefactor* pipeline was run through a single data set on three clusters and a personal machine. A number of insights were made using the high resolution timing data collected from this package (such as in Figure 4.10) and are listed in Table 4.1. In the future, we'll apply the *pipeline\_collector* software to the more complex LOFAR DD pipeline, *ddf-pipeline*<sup>12</sup>.

The slowest processing steps for the *prefactor* pipeline were identified as the *calib\_cal* and *gsmcal\_solve* steps. While the data can fit into the RAM for all of the processing machines, it is much larger than the processor's internal cache (Figure 4.6). The discoveries made concerned the memory hierarchy in Figure 4.6. Results labeled **R2**, **R8** and **R9** related to the CPU performance; **R2**, **R6** and **R7** related to the Cache performance; **R3** and **R5** related to the Memory usage and **R4** discussed the Disk speed.

Faster processors did not accelerate the *gsmcal\_solve* step significantly, as this step streams data between the RAM and CPU. As the CPU speed increases, streaming applications become bottlenecked by the throughput of data into the CPU

<sup>12</sup><https://github.com/mhardcastle/ddf-pipeline>

from RAM. As the *gsmcal\_solve* algorithm iteratively calibrates chunks of the data, these chunks need to be loaded from disk once, however they are moved from RAM to CPU multiple times during calibration.

Similarly, the *calib\_cal* step is more dependent on memory throughput than on CPU speed as this step moves data to and from memory frequently. This step also does minimization looping over the data set. As the data set does not fit in the cache, parts of it need to be constantly moving from memory and back. Figure 4.8a shows that the machine with the smallest LLC cache runs the *calib\_cal* step the fastest. This is likely a combination of the benefit of faster RAM and poor cache optimisation for this software. The same effect is much less pronounced in Figure 4.8b, suggesting that software optimisation at least plays a part in the outliers for the laptop machine.

### 4.5.1 Recommendations

Based on these results, the top hardware recommendation is that *prefactor*'s slowest steps can be accelerated by running on machines with faster memory or upgrading the memory of the current machines. The two slowest *prefactor* steps showed improvements on machines with faster RAM.

One software recommendation is to improve the efficiency of the *calib\_cal* step through refactoring or by replacing the software package used. Unfortunately, the software used for the *gsmcal\_solve* step cannot be used for the *calib\_cal* step as it is not yet able to correct for Faraday Rotation [105], making it impossible to currently use the software used by the *gsmcal\_solve* step. Faraday Rotation has recently been implemented in a development version of the *prefactor* pipeline and is currently undergoing testing. This version of the pipeline will be implemented by September 2018.

Additionally, the large number of data cache misses recorded for the *calib\_cal* step suggests that its source code is not optimised for multi-threaded processing. Data cache misses are often encountered when multiple threads have instructions on the same cache line<sup>13</sup>, forcing the memory controller to move this cache line between cores [109]. This can also explain the large number of stalled cycles (Fig. 4.13a) and low number of full issue cycles (Fig. 4.13b) for the *calib\_cal* step. It is recommended to further study the inefficiencies of *calib\_cal* or to replace it with a newer software. If the software processing for this step is updated, analyzing

<sup>13</sup>A cache line is a row of cache memory which is loaded into CPU as a single unit [112]

the cache and CPU performance of the new software will be necessary to determine whether it efficiently uses the available computational resources.

Finally, we discovered that compiling the software on a virtual machine did not lead to a processing slowdown. This means that the current slowest *prefactor* steps are not optimised to use advanced processor instructions. Nevertheless, the resulting cross-compatibility is an encouraging result as it will allow to easily distribute pre-compiled versions of the software without increasing the processing time. We recommend continuing CVMFS deployment of LOFAR software.

## 4.6 Conclusions

In this chapter, we present a novel system for automated collection of performance data for complex software pipelines. We use this suite to study the LOFAR *prefactor* pipeline. The results are discussed aiming to understand the effect of different hardware parameters on the data processing. To do so, we run the pipeline on four different machines.

The software automatically collects performance data at the operating system level without impacting processing time. Data for each pipeline step is extracted using the OpenTSDB API, plotted and analyzed. Additionally, the *pipeline\_collector* suite is easy to extend with new collectors that record more detailed time-series data for each pipeline step. The performance data is stored in the time series database OpenTSDB.

Here, we used this data to find 9 insights into the LOFAR *prefactor* pipeline listed in Table 4.1. The implementation details are described in 4.A.

The *prefactor* pipeline is used to do the initial processing for over 3000 observations that are part of the LOFAR SKSP Tier 1 survey. However, this pipeline is also used for lots of other LOFAR data sets outside the SKSP project. We have shown that increasing the RAM throughput is the easiest way to speedup *prefactor* processing. Running the *calib\_cal* step on hardware with RAM faster than 4 GB/s will save up to 700k CPU hours for the 3000+ unprocessed data sets. This throughput increase will also speedup the *gsmcal\_solve* step by 30% saving an additional 400k CPU hours. This is a significant fraction of the estimated 2,400k CPU hours required to process this data with the *prefactor* pipeline.

As shown in this work, we can correlate the performance of the LOFAR software with different hardware specifications. Additionally, the data sets can vary

in size and job overheads on the compute cluster can depend on the processing parameters. All of these parameters affect the processing latency for the calibration and imaging pipelines. As such, a thorough parametric model is required to further optimise the end-to-end LOFAR processing pipeline and predict processing times on future clusters.

The design of this utility makes it easy to apply to future LOFAR pipelines. It is important to note that *pipeline\_collector* is general enough that it can be used by other scientific pipelines, with no modification of the pipeline. Integrating *pipeline\_collector* with a different pipeline requires only minor work. In future work, we will integrate *pipeline\_collector* with *ddf-pipeline*. We will use this data to create a performance model of the full LOFAR imaging pipeline [Ivanweeren2016, 17, 83], including the DI step, implemented by *prefactor*, and the DD step, implemented by *ddf-pipeline*. This model will make it possible to first understand and then optimise the LOFAR pipeline and suggest for hardware and software improvements.

## 4

## 4.A Performance Collection Implementation Details

In this work, we have developed the *pipeline\_collector* suite, aimed at collecting detailed time-series information from distributed scientific pipelines.

The *tcollector* package is a python software suite that can collect system performance data at predetermined intervals. The package is designed to monitor the performance statistics for web-servers and cluster nodes. The *tcollector* software records time series of the different performance metrics and sends them to a Time Series Database through HTTP. The Time Series Database, OpenTSDB stores the data in an HBase [113] instance at the performance collection server. Users interested in plotting time series can plot real time or historical data through an HTTP interface with OpenTSDB. With a central performance collection server, data from multiple processing sites can be collected and analyzed.

Tcollector formats the time series information in four fields. First is the name of the metric which is measured. Second is the UNIX timestamp. Third is the time series recorded as an integer or a float. Finally, a set of tags (key-value pairs) can be added to the data point. These four fields are discussed below and can be seen on the right side of Figure 4.14.

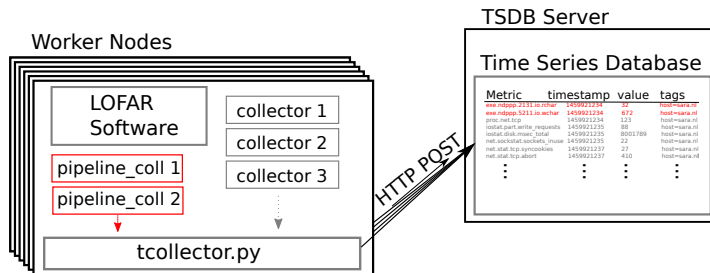


Figure 4.14: Communication between worker nodes and the TSDB server, including the `pipeline_collector` modules (in red). The `pipeline_collector` suite collects information on the running LOFAR pipelines, while the rest of the `tcollector` package collects system performance data. The existing `tcollector` package and its collectors are shown in gray. The collectors in gray only record metrics from the global system.

#### 4.A.1 The `pipeline_collector` suite

The `tcollector` package cannot collect data on individual processes, nor can it associate these processes with specific steps of a data processing pipeline. We’ve supplemented the software with the `pipeline_collector` suite<sup>14</sup> using an executable that monitors a pipeline’s running processes<sup>15</sup>. When an executable that is part of the LOFAR pipeline launches, a dedicated collector begins reporting information on the individual process. Every time a new processing step starts, the `prefactor` pipeline records the step name in a log file. `pipeline_collector` determines the current running step using this log. Running the LOFAR processing concurrently with the `tcollector` package gives us per-step performance data without changing or slowing down the LOFAR `prefactor` pipeline. Furthermore, `pipeline_collector` can be integrated with any processing pipeline as long as each pipeline step’s name is recorded at its launch.

The `pipeline_collector` suite sends data to the time series database in the same format as the rest of the collectors included in the `tcollector` package.

#### 4.A.2 Setting up for future pipelines

The setup options for `pipeline_collector` are stored in a configuration file in the root directory of the package. This file holds the sample interval, executables to monitor and the location where `pipeline_collector` can read the current pipeline step

<sup>14</sup> Located at [https://gitlab.com/apmechev/procfs\\_tcollector.git](https://gitlab.com/apmechev/procfs_tcollector.git)

<sup>15</sup> Located at <https://github.com/apmechev/procfsamp>

The *pipeline\_collector* suite reads the current pipeline step from a file, the location of which is specified in the configuration. This file needs to be updated each time the pipeline begins a new step. For LOFAR we have a script running with the pipeline, and determining the current step using the pipeline logs. As each pipeline has a unique sequence of steps, the current step needs to be recorded in a file in order for *pipeline\_collector* to report it to the time series database. The location of the file recording the current pipeline step is read from the configuration file.

Next, the names of the specific processes need to be included in the configuration file. In the case of LOFAR, we select the NDPPP, bbs-reducer and losoto processes. The *pipeline\_collector* searches the running processes for the current user for these process names and launches a collector for each new process launched by the current step.

**Original Abstract:**

The LOFAR radio telescope creates petabytes of data per year. This data is important for many scientific projects. The data needs to be efficiently processed within the time span of these projects in order to maximize the scientific impact. We present a workflow orchestration system that integrates LOFAR processing with a distributed computing platform. The system is named Automated Grid-enabled LOFAR Workflows (AGLOW). AGLOW makes it fast and easy to develop, test and deploy complex LOFAR workflows, and to accelerate them on a distributed cluster architecture. AGLOW reduces the setup time of complex workflows: typically, from months to days. We lay out two case studies that process the data from the LOFAR Surveys Key Science Project. We have implemented these into the AGLOW environment. We also describe the capabilities of AGLOW, paving the way for use by other LOFAR science cases. In the future, AGLOW will automatically produce multiple science products from a single data set, serving several of the LOFAR Key Science Projects.

This chapter is based on the work of **A.P. Mechev**, J.B.R. Oonk, et al. “Fast and Reproducible LOFAR Workflows with AGLOW”. in: *2018 IEEE 14th International Conference on e-Science (e-Science)*. Oct. 2018, arXiv:1808.10735. DOI: 10.1109/eScience.2018.00029. arXiv: 1808.10735 [astro-ph.IM]

## 5.1 Introduction

Data sets in radio astronomy have increased 1000-fold over the past decade[114]. It is no longer feasible to move, store and process these data sizes at university clusters, nor to process these data manually. LOFAR, the Low-Frequency Array[77] is a modern and powerful radio telescope that creates more than 5 Petabytes of data per year. At present, the majority of LOFAR time is allocated to several Key Science Projects (KSPs)[25]. These projects need to process hundreds or thousands of observations. Typical observations produce approximately 14 TB of archived data. Obtaining high fidelity images from this data requires complex processing steps. To manage and automate the data processing, workflow management software is needed. This software needs to accelerate LOFAR processing on a High Throughput Computing (HTC) cluster while ensuring it is easy to prototype, test, and integrate future algorithms and pipelines.

To automate LOFAR data processing, we have worked with the LOFAR Surveys KSP (SKSP). Together, we designed a software suite that integrates LOFAR software[95] with the Dutch Grid infrastructure[115]. This software, based on Apache Airflow<sup>1</sup>, makes it easy to add future science cases, extend and modify pipelines, include data quality checks, and rapidly prototype complex pipelines. For the SKSP use cases, AGLOW achieves a significant reduction in development time: from months to days, allowing researchers to concentrate on data analysis rather than management of processing. Additionally, and perhaps more importantly, the software versions and repositories used are defined within the workflow. This makes reproducibility an integral part of the AGLOW software. Finally, the software is built to leverage an HTC cluster by seamlessly submitting the processing jobs through the cluster's job submission system[116]. The work presented here builds on our previous work parallelizing single LOFAR jobs[35] on a distributed environment. The majority of processing was done at SURFsara at the Amsterdam Science Park[85], which is one of the sites used by the LOFAR Long Term Archive (LTA)<sup>2</sup>. Ongoing efforts include scheduling and processing data at clusters in Poznań in Poland and Jülich in Germany.

**Contributions:** The main features of the AGLOW software are the following:

- Integration of the Grid middleware with Apache Airflow, allowing us to dy-

---

<sup>1</sup><https://airflow.apache.org/>

<sup>2</sup><https://lta.lofar.eu>

namically define, create, submit and monitor jobs on the Dutch national e-infrastructure.

- Integration of the LOFAR (LTA) utilities in Airflow, facilitating pipeline developers to automate staging (moving from tape to disk) and retrieval of LOFAR data.
- Integration of the SURFsara storage with Airflow, making LOFAR pipelines aware of the storage layer available at the Dutch national e-infrastructure.
- Ease of creating simple software blocks, with which users can integrate and test their pipelines.
- Storing all software versions and script repositories as part of the workflow to make LOFAR processing reproducible and portable.

**Outline:** The organisation of this manuscript is as follows: We provide background on data processing in radio astronomy and why LOFAR science requires complex workflows and cover workflow management algorithms and capabilities (section 5.2). We discuss related work in workflow management (section 5.3). In section 5.4, we introduce our software and two use cases. Both of our use cases require acceleration at an HTC cluster and automation by a workflow orchestration software. We follow these examples with details on the integration between LOFAR software, LOFAR data and the resources at SURFsara in Amsterdam in section 5.4.2. Finally, we discuss our results (sect. 5.5) and look ahead to the demands of future LOFAR projects and upcoming telescopes in section 5.6.

## 5.2 Background

This work lies at the intersection of Radio Astronomy and Computer Science. The goal of the study is to leverage the flexibility of an industry standard workflow management software and use CERN's Worldwide Computing Grid<sup>3</sup> at SURFsara [117] to accelerate reproducible processing of LOFAR data.

A single LOFAR surveys observation is recorded in distinct frequency chunks (henceforth called 'Subbands'), each of which is uploaded to the LTA as a separate file. Some of the processing steps require the entire frequency information, while other steps can run independently and operate on a single Subband. The

---

<sup>3</sup><http://wlcg.web.cern.ch/>

latter steps can be easily accelerated on an HTC cluster by taking advantage of the data level parallelism.

Multiple scientific projects may desire to run different processing steps on a single LOFAR observation. To minimize time spent on retrieving data from the LTA and eliminate re-processing of data, pipelines for multiple science cases need to be integrated together. This integration should be done by a software that encodes the dependencies between different steps and automatically executes processing steps once their dependencies have been met. Software packages that solve these challenges are called ‘workflow management software’ (see, e.g., [118–120].)

### 5.3 Related Work

A workflow is described by a set of tasks. The dependencies between these tasks are encoded in a Directed Acyclic Graph (DAG) [121]. This data structure imposes a strict dependency hierarchy between the tasks [122]. This means that there exists a well-defined execution order and a well-defined list of dependencies for each task. The execution order is typically determined by algorithms such as Kahn’s algorithm [123] or a depth-first search [124].

Workflow management software is used in various fields from research to industry. In biology, gene sequencing and analysis pipelines require automation of multiple processing steps. In gene sequencing, Toil<sup>4</sup> has been successfully used to automate RNA sequence analysis[125]. Additionally, many software teams in biotech develop their own in-house workflow management software [126].

Currently, we can parallelize a single processing step of the pipeline using the Grid LOFAR Tools (GRID\_LRT<sup>5</sup>) [35]. The LOFAR Surveys science cases incorporate multiple steps with inter-linked dependencies. Resolving these dependencies can be done efficiently by a comprehensive workflow orchestration software. The purpose of such software is to resolve dependencies between the multiple tasks in a workflow, execute these tasks, and track the status, logs, output, and run time of each task.

In astronomy, workflow systems have been developed that are telescope specific, such as ESOReflex[127] by the European Southern Observatory. Other

<sup>4</sup><https://toil.readthedocs.io>

<sup>5</sup>[https://github.com/apmehv/GRID\\_LRT](https://github.com/apmehv/GRID_LRT)

projects, such as *astrogrid*<sup>6</sup> and 'Workflow 4Ever'<sup>7</sup>, have either been completed or are no longer supported. The *astrogrid* project, for example, was a collaboration to create standards, infrastructure, and software for distributed astronomical processing. Its operation phase spanned 2008-2010. *Workflow4Ever*, likewise, has been out of support since 2013. To ensure continuing support for the LOFAR workflows, we have decided to use a leading enterprise workflow management software, *Airflow*.

*Airflow* is an open source Python software package developed by Airbnb<sup>8</sup> to manage complex workflows. It encodes workflows in Python files and makes it easy to re-use, re-arrange, schedule and execute blocks in a user-defined workflow. *Airflow* is capable of scheduling and executing workflows by resolving the dependencies between tasks and scheduling these tasks for execution. The software uses a metadata database<sup>9</sup> to retain metadata such as task state, execution date, and output. While *Airflow* allows building workflows easily from Python and bash functions, it can easily be extended to support custom processing scenarios. Additionally, *Airflow* conforms to the Common Workflow Language (CWL) [128] standard using the *cwl-airflow* package [129], meaning it can execute CWL workflows as well. Finally, *Airflow* is part of the Apache incubator and upon certification will receive continual support by the Apache software foundation<sup>10</sup>.

## 5.4 AGLOW

Complex astronomical pipelines are time consuming to develop and operate. Furthermore, they may evolve rapidly to incorporate new processing techniques or requirements. Migrating these pipelines to a distributed, high throughput environment is often justified, or even required, in order to meet the timelines set by scientific projects. The time saved by running on a cluster must be balanced by the flexibility and development time required to implement or update the scientific pipelines. To address these concerns, we have developed a software package, Automated Grid-enabled LOFAR Workflows (AGLOW)<sup>11</sup>. AGLOW is based on *Airflow* and the LOFAR software and addresses issues with automation and acceleration of LOFAR processing.

---

<sup>6</sup><http://www.astrogrid.org>

<sup>7</sup><http://wf4ever.github.io/ro/>

<sup>8</sup><https://www.airbnb.com/>

<sup>9</sup>In our case implemented by PostgreSQL

<sup>10</sup><https://www.apache.org/>

<sup>11</sup><https://github.com/apmechev/AGLOW>

With AGLOW, we can translate LOFAR pipelines into DAGs. We provide the tools that enable users to easily implement LOFAR science pipelines and execute them on a distributed architecture. Using these tools, the data processing required by various LOFAR science cases is automated and accelerated.

### 5.4.1 AGLOW: Case Study

For our case-study, we have chosen two ways to process LOFAR Surveys data: coverage and depth. The Surveys Key Science Project (SKSP)[25] is an ambitious project to map the northern sky at low frequencies using the Dutch LOFAR stations. These maps will help understand the formation and evolution of massive black holes, galaxies, clusters of galaxies and large-scale structure of the Universe.

The LOFAR surveys observations consist of several tiers with the widest Tier (Tier 1) covering the whole sky visible from the Northern Hemisphere with 3168 observations of 8 hours each [25]. The other tiers (Tier 2, Tier 3) consist of much longer observations of smaller sections of the sky and can collect hundreds of hours of data for a single direction. The deepest single field being analyzed, in collaboration with the EoR group, is the North Celestial Pole (NCP) field which has  $\sim 1700$  hrs of observations to date. Processing this data will create an image with an unprecedented resolution and sensitivity. Here we have implemented processing pipelines for both the Tier 1 data and Tier 2 data into AGLOW.

The scientific importance of these two examples, as well as the large processing requirements, make them ideal candidates for acceleration and automation with AGLOW.

#### Surveys Project: All Sky Survey

The main driver for the development of AGLOW and its constituent packages has been the LOFAR SKSP Project. A typical 8-hour observation produces 14 TB of data. This data is eventually reduced to several hundred gigabytes. Data needs to be processed by two pipelines: first by the Direction Independent (DI) pipeline, and then by the Direction Dependent (DD) pipeline.

We have split the DI pipeline into four stages, and the DD pipeline into two subsequent stages. Splitting up the pipelines in stages allows speedup through parallelization for the stages that can benefit from data-level parallelism. Additionally, this setup allows fault tolerance and easy re-processing. Importantly, with AGLOW,

adding new functionality to the pipeline is easy and can be done at any time without disrupting current processing. Current SKSP processing is easily started by launching a new DAG-run in Airflow. We schedule these runs daily at midnight, however they can also be started manually by users.

### Deeper Surveys Fields

To create deep images of a single field, minor modifications were made to the processing pipeline described in the previous section. Scripts were included to re-align the data in the frequency axis, and the DD processing steps include an extra final combination step that stacks multiple observations. Being able to rapidly test alternative processing strategies is crucial to creating a deep image of the NCP field. With the success of this project, future deep LOFAR observations will be processed with these pipelines.

#### 5.4.2 AGLOW: Implementation

AGLOW combines the LOFAR software, the Grid LOFAR Tools (GRID\_LRT), and Airflow to allow automation and makes large-scale LOFAR processing easily reproducible. The components of the AGLOW software are shown in Fig. 5.1. In this section, we will discuss these components and their functions.

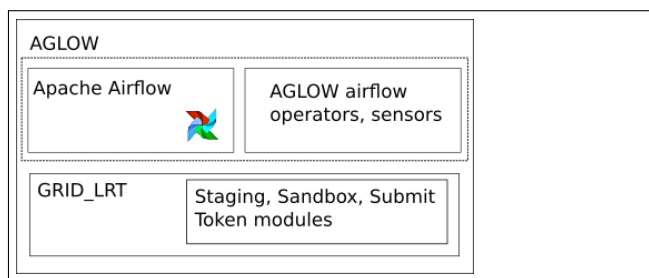


Figure 5.1: Design of the AGLOW software, incorporating Airflow, the GRID\_LRT package[35] and custom operators designed to integrate LOFAR software, Grid middleware and dCache storage. GRID\_LRT is a software package developed to parallelize single LOFAR jobs at SURFsara. It contains several modules to help set up, and launch jobs on an HTC cluster at SURFsara. Airflow is a stand-alone package by Airbnb, which is extended with several classes that couple Airflow with the Grid infrastructure. These classes are collectively named the AGLOW operators/sensors.

## GRID LOFAR Tools and LOFAR software

We have previously developed tools to create LOFAR jobs and launch them on a distributed infrastructure[35]. These tools have matured to a point where it is easy to both plug and play existing scripts and extend the framework to add more complex pipelines. These steps make it possible for a user to batch execute bash or Python scripts on their LOFAR data in parallel. After the scripts are executed, the results are uploaded to shared dCache storage[130] at SURFsara[85].

More complex steps use additional Github repositories, such as the prefactor<sup>12</sup> for direction independent calibration or DDF-pipeline<sup>13</sup> for direction-dependent calibration and imaging. The sequence of steps is encoded in parameter-set files (parsets), which can be modified and dropped into AGLOW depending on the processing requirements. Users have full control over which git repository, branch and commit number to use for each processing step.

With AGLOW, we can easily include the DDF-pipeline and prefactor repositories, as well as any other scripts. Since these scripts are tracked by git [131], a full commit and branch history of the scripts is available. We use this history to make processing reproducible, by using the same git-commit for all LOFAR data sets.

In addition to these script repositories, we have integrated the most common software packages used to process LOFAR data with AGLOW. These are the Default Pre-Processing Pipeline (DPPP)[95], the LOFAR Solutions Tool (LoSoTo), WSclean [132], AOflagger[133], CASA[134], pyBDSF[135], DDFacet[19] and KillMS[83, 136].

## Extending Airflow

Two types of modifications were made to Airflow to allow processing on a Grid environment. First, functions were added to check the number of files located in intermediate grid storage. We use this to decide whether to stage files, or whether enough files have been successfully processed by a previous task.

Second, more complex tasks were implemented as Airflow operators or sensors (Figure 5.2). These tasks include creating job description files, setting up

<sup>12</sup><https://github.com/lofar-astron/prefactor>

<sup>13</sup><https://github.com/mhardcastle/ddf-pipeline>. DDF-pipeline is a leading example of a Direction Dependent calibration pipeline used for LOFAR data. It uses DDFacet [19], KillMS [83] and to create high quality images.

job scripts, launching jobs using the gLite workload management system and monitoring the status of these jobs. Future additions will include operators that evaluate the current cluster workload and make decisions on location to launch the data processing. With the AGLOW package, such tasks are easy to implement without modifying or interrupting processing. This leads to an easily reproducible, intelligent scientific processing that is also efficiently executed and requires minimal interaction. The operators and sensors added to Airflow are shown in Fig. 5.2. The most commonly used ones are described below:

**Staging** AGLOW stages data through the ASTRON staging API, which is made available through the GRID\_LRT software <sup>14</sup>. This API loads the user's LTA credentials from a file. Using these credentials, AGLOW stages the required files at the respective LTA site. The AGLOW staging operator waits until all the files have been staged before returning a successful exit status.

**Sandbox** Scripts to retrieve, process and upload the data are stored in a sandbox, as described in [35]. The LRT\_Sandbox operator that builds this sandbox directly uses the GRID\_LRT Sandbox module to build and upload the sandbox according to the user-provided configuration file.

**Token Creation** Jobs parallelized with GRID\_LRT store processing parameters inside CouchDB documents called PiCaS Tokens. Each token corresponds to one grid processing job. These tokens are created by the LRT\_Token operator in AGLOW. This operator uses the GRID\_LRT Token module to define tokens and create them from a list of links to LOFAR data.

**Job Submission** The job submission at SURFsara is done through the GRID\_LRT submit module. This module is a python interface to the glite job submission system. It is wrapped in the glite-submit operator in AGLOW, which returns the ID number of the job. AGLOW includes a gliteSensor, which uses this ID to wait until all the processing jobs have finished.

Using AGLOW to accelerate the execution of a pipeline requires deciding how to split the processing to benefit from parallelization. Once the steps to be parallelized are selected, users can add git repositories of scripts to the configuration file. Next, each step is added to a Python script called the DAG file. This file is placed

<sup>14</sup><https://grid-lrt.readthedocs.io/en/latest/staging.html#grid-lrt-staging-stage-all-lta>

in the Airflow's dags folder, which adds it to AGLOW. An example DAG file can be found at <https://aglow.readthedocs.io/en/latest/dags.html>. Users who want to implement a new workflow design the DAG files and add them to Airflow's dags folder. Workflows are migrated to a new server, by transferring the DAG and configuration files to the new AGLOW instance.

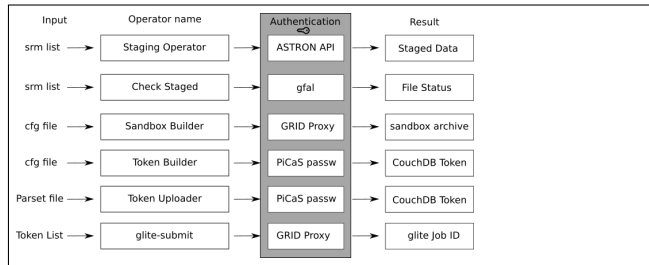


Figure 5.2: Airflow Operators for Staging LOFAR data, creating job descriptions and submitting jobs to the Dutch grid. On the left is the input given to each operator. ‘Srm lists’ are lists of links to data at the LOFAR LTA or located on the SURFsara dCache storage. Parsets are files specific to ‘prefactor’ and ‘DDF-pipeline’ and define the processing for each pipeline step. Finally, the ‘Sandbox’ and ‘Token’ operators read their parameters from a configuration file. The use of a scripts sandbox and job description tokens is detailed in our previous work[35]. Documentation of the operators is available at <https://aglow.readthedocs.io/en/latest/operators.html>

### 5.4.3 AGLOW: Jobs

Once LOFAR observations are downloaded from the LTA, they are typically processed with several packages before producing a science ready data set. We have integrated these packages with Airflow to make it easy to create complex LOFAR workflows.

Each of the processing steps above requires extra set-up to process on the Dutch Grid infrastructure. The job scripts setup, job description, and job submission are done by the GRID\_LRT package[35]. With AGLOW, we automate this setup, enabling users to focus on developing more comprehensive data processing pipelines. Below we outline several possible steps a user can use in their pipeline.

#### DPPP Parset

The DPPP software is used extensively in LOFAR data processing. It has many capabilities such as flagging bad data, averaging data in time and frequency, and

calibrating the data with a sky-model.

The input parameters of this software are stored in a text file called a parset. The input data and the DPPP parset are sufficient to define a DPPP execution step. As noted in section 5.2, LOFAR data is split in frequency into Subbands. Much of the DPPP processing, such as averaging and flagging, can be done independently for each Subband, thus they can be processed on independent machines. This parallelization makes these steps a perfect candidate for an HTC cluster. For a data set that is split into 244 Subbands, 244 jobs are launched concurrently.

In Airflow, the DPPP parset task is encoded in a DAG (Fig. 5.3). The DPPP DAG is a linear workflow that consists of the 'sandbox' setup, creation of the job-description documents, uploading of the DPPP parset and job launching and monitoring.

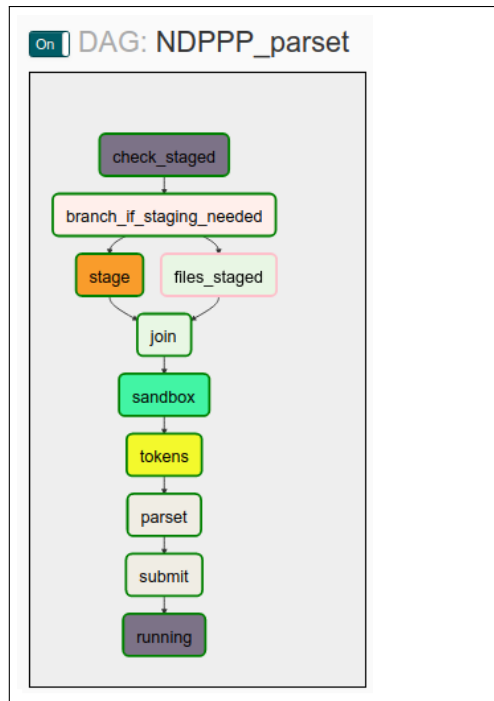


Figure 5.3: Render of the DPPP parset DAG in the Airflow User Interface. This view shows the setup and submission steps. Even this simple DAG can include branching options such as the `branch_if_staging_needed` task which checks if the data is not staged and stages it. All of the operators in this figure are part of the AGLOW software. Their inputs and outputs are shown in Fig. 5.2. Using configuration files, the NDPPP DAG can be used by different users for different science cases making it portable and maintainable. These features make reproducible science with LOFAR data easy.

## WSclean Job

The WSclean[132] package is used to create an image from a LOFAR data set. This software has a very wide range of parameters options, however, it cannot take a parset file as an input. Instead, the parameters are specified in the command line. In the AGLOW implementation, we parse all the command-line parameters from a text file, referred to as the 'wsclean parset'. This file is added to the jobs in the same way as the DPPP parset, i.e. using the *Token Uploader* Operator. The DAG for the wsclean software uses the same blocks as the DPPP DAG, with different configuration and parset files. The reuse of Airflow operators makes maintainability of all tasks easier.

### 5.4.4 Shell/Python Script

Users that require the run of multiple software packages on a single data set can craft a custom shell or Python script that executes these steps using the LOFAR tools during a single distributed job. This option increases flexibility and minimizes the overhead associated with scheduling and running multiple jobs in sequence. At the workflow orchestration level, we use the same Airflow operators as the above tasks. The script is uploaded to the job description database using the *Token Uploader* Operator. It is executed once the jobs are launched.

Currently only the LOFAR Spectroscopy project uses custom shell scripts to process LOFAR data. A recent study of carbon recombination lines used a custom bash script to calibrate and image LOFAR data on the SURFsara GINA<sup>15</sup> cluster [137].

### 5.4.5 Prefactor parset

The input to the prefactor pipeline software is a parset file which describes a linear workflow. The description of this workflow consists of a list of processing steps and their associated parameters. The 'prefactor' package uses the LOFAR software to do the direction-independent calibration of the archived LOFAR datasets. Prefactor steps are executed by the generic pipeline framework[95]. While this framework can run a sequential pipeline, it is not capable of conditional branching nor parallelization on all cluster architectures. The original goal of the GRID\_LRT software

<sup>15</sup>The GINA cluster is an HTC cluster located at SURFsara integrated with the Dutch Grid initiative. It supports massively parallel processing which is required to efficiently process LOFAR data with *prefactor*.

was to tackle the parallelization challenge while AGLOW solves the additional challenge of pipeline management.

We have already processed more than 50 datasets through the ‘prefactor’ DAG using AGLOW. The full ‘prefactor’ pipeline is shown in figure 5.4. This DAG shows the four processing steps as well as additional Python operators that manage the staging and result verification.

### 5.4.6 DDF-pipeline

The final AGLOW DAG is the implementation of the DDF-pipeline repository which is a pipeline that is extensively used by the LOFAR surveys KSP and is described in detail in [25]. This pipeline operates on the products of the prefactor pipeline and consists of a series of calibration and imaging loops with the objective of creating a final science quality image. For each of these loops the majority of the processing time is spent in DDFacet[19] and KillMS[83, 136] steps that perform the direction-dependent imaging and calibration respectively.

In total, DDF-pipeline takes  $\sim 4$  days of processing to complete. As DDF-pipeline creates large intermediate files we have so far not divided the pipeline into too many steps to avoid filling the storage on the GINA cluster. However, we have split the pipeline into two steps and there is further potential for parallelization that will be implemented in the future.

### 5.4.7 Linking Multiple Jobs

Pre-processing of LOFAR SKSP data can be done by a single DPPP task, with 244 jobs running in parallel. More complex LOFAR pipelines will include multiple processing tasks as well as tasks responsible for job setup. Therefore, it is important to facilitate running multi-step pipelines with AGLOW.

Creating workflows by defining dependencies between tasks is a core Airflow capability. We use this functionality to link multiple steps of a LOFAR pipeline together. In the SKSP pipeline, we take advantage of the data level parallelism for the initial processing steps for the calibrator and target. The other two steps are run as a single grid job. Switching the parallelization for each step is done by changing the number of datasets per node parameter in the configuration file for each step.

## 5.5 Results and Discussions

The implementation of AGLOW makes it possible to efficiently process LOFAR data with minimal user interaction. The scheduling algorithm automatically launches pipelines, meaning that there is little time spent between runs. Additionally, controlling/fixing the version of the scripts is done by specifying the commit of each script repository. This makes data processing easily reproducible. Once the dependencies of multiple science pipelines have been encoded in a DAG, Airflow efficiently executes this DAG, running tasks in parallel where possible.

An important feature of AGLOW is the loose coupling between pipeline logic, software versions, pipeline parameters, and datasets. The goal of this decoupling is to give users complete control over all the processing variables. With AGLOW, one can develop the pipeline logic independently of the LOFAR software versions and conversely update the LOFAR software and script repositories independently from the pipeline logic. Finally, the Airflow operators are themselves decoupled from the scientific pipelines. As these operators are reused, this decoupling makes them easy to maintain and extend.

The first LOFAR processing pipeline integrated with AGLOW was a single linear workflow, with only one submission to the compute cluster. This workflow is used to reduce the data size making data retrieval to research institutes less time consuming. We offer this workflow as a service to LOFAR users who do not have a high-bandwidth connection to the LOFAR Archive.

A more complex pipeline was implemented: the LOFAR direction independent calibration pipeline (‘prefactor’). The scientific importance and complexity of this pipeline make it a good case study for the capabilities of the AGLOW software. We show that AGLOW’s design allows integration of more complex data processing workflows with the Dutch Grid resources. These workflows can be either used by PIs of LOFAR projects or offered as a processing service to the wider astronomical community. Since March 2018, more than 300 ‘pre-factor’ datasets have been processed with the SKSP workflow. With AGLOW, we have been able to process data with an average cadence of 50 observations per month.

In large part thanks to their flexibility, automation, and Grid integration, AGLOW and GRID\_LRT have become a standard part of the Direction Independent processing for the LOFAR SKSP project.

## 5.6 Conclusions

In this work, we have detailed a comprehensive workflow management software for processing radio astronomy data on a distributed infrastructure. We leverage an industry standard workflow management software, Airflow. Using its capabilities, we make it possible to build, test, automate and deploy LOFAR pipelines on short timescales, generally from months to days. With the flexibility of Airflow's Python and Bash operators, users can design their own workflows, as well as co-ordinate more complex science cases. In this way, AGLOW facilitates reproducible processing of scientific data. In the future, AGLOW will support additional LOFAR science cases including Long Baselines and Spectroscopy. In this article, we have described our implementation of the data processing pipelines used by the LOFAR Surveys Key Science Project.

Future work includes further de-coupling of the Grid-setup and pipeline logic. We will do this by creating 'sub-dags' (details in 5.6.1) for each type of LOFAR jobs. Using these sub-dags will reduce the complexity of scientific workflows while also making the code even more reusable and thus easier to maintain and upgrade. Efforts to integrate processing at the other two LTA sites, (Jülich and Poznań) have already started with 'prefactor' runs being performed on Jülich using a modified version of the SKSP workflow. The software also currently works on the Eagle cluster at Poznań. Combining the Jülich and SURFsara workflows will be done in the future so that AGLOW can track and start processing at multiple clusters.

Finally, AGLOW can be used as a 'LOFAR As A Service' model. In this model, users only provide an observation ID and processing parameters and receive the final results upon job completion. This model will build upon previous success offering LOFAR processing to users without login to the GINA cluster [138]. This previous work was already useful for studying radio absorption in Cassiopeia A [29] and a 'data-to-images' service will be valuable to the whole LOFAR community.

Our experience with automating LOFAR scientific workflows on a distributed architecture will be valuable when setting up data processing for future Radio Telescopes such as the Square Kilometer Array [139].

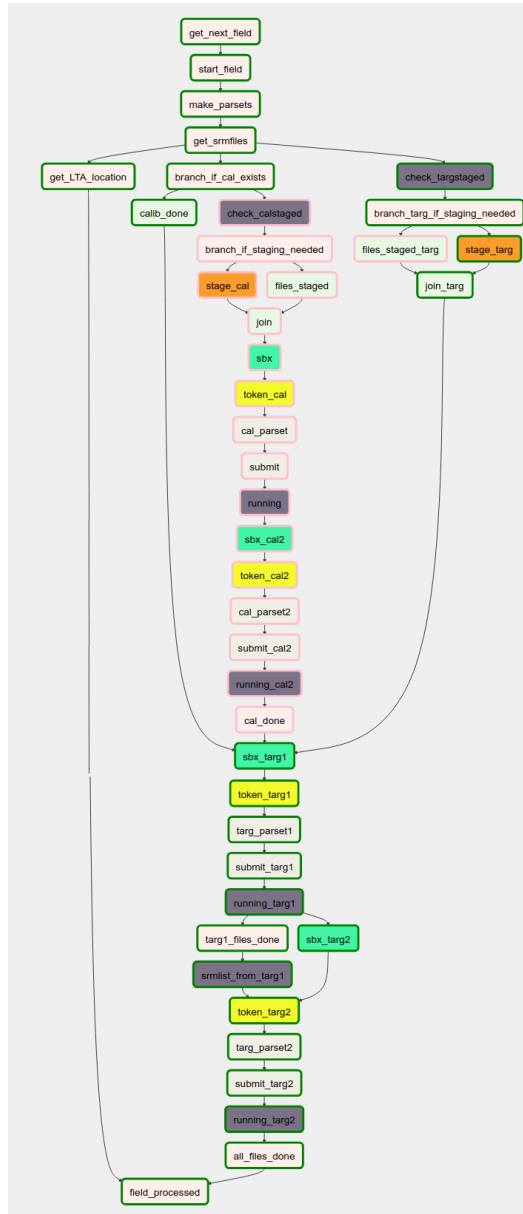


Figure 5.4: Workflow for the prefactor pipeline. Here we show the reuse of AGLOW operators for the four prefactor steps. In addition to the LOFAR processing, we also have conditional operators to skip processing of the calibrator if it has been previously processed. This is done by the ‘branch\_if\_cal\_exists’ task. We also have a final step that checks if all the results have been uploaded, done by the ‘all\_files\_done’ task. Likewise, quality checks can be added in this workflow wherever needed.

## APPENDIX

The LOFAR SKSP workflow is shown in Figure 5.4. This figure shows how reuse of the staging, setup operators, and glite-wms sensors makes maintainability easy and allows rapid prototyping of complex pipelines.

This workflow additionally takes advantage of Airflow's *PythonOperator* to check if the LOFAR data is on disk at the archive and whether all final products were uploaded by each step. AGLOW also allows for staging the calibrator and target files concurrently. When the data is staged, Airflow continues with the processing of that data.

### 5.6.1 Sub-DAG

Airflow allows developers to include entire DAGs as a single task in their workflow. Airflow can trigger a DAG execution based on parameters provided by the parent DAG. This feature makes it possible to concatenate short, commonly used tasks into DAGs and call them in a parent workflow. Using sub-DAGS makes the code more maintainable and easy to use, while it makes workflows simpler. For LOFAR, Sub-DAGs are used to automate job submission, making the resulting scientific workflows simpler.



**Original Abstract:**

LOFAR is a leading aperture synthesis telescope operated by the Netherlands with stations across Europe. The LOFAR Two-meter Sky Survey (LoTSS) will produce more than 3000 14 TB data sets mapping the entire northern sky at low frequencies. All of the data produced by LoTSS needs to be processed by the LOFAR Direction Independent (DI) pipeline, `prefactor`. Understanding the performance of this pipeline is important when trying to optimise the throughput for its large projects, such as LoTSS and other deep surveys. Making a model of its completion time will enable us to predict the time taken to process large data sets, optimise our parameter choices, help schedule other LOFAR processing services, and predict processing time for future large radio telescopes. We tested the `prefactor` pipeline by scaling several parameters, notably number of CPU's, data size and size of calibration sky model. We present these results as a comprehensive model which will be used to predict processing time for a wide range of processing parameters. We also discover that smaller calibration models lead to significantly faster calibration times, while the calibration results do not significantly degrade in quality. Finally, we validate the model and compare predictions with production runs from the past six months, quantifying the performance penalties incurred by processing on a shared cluster. We conclude by noting the utility of the results and model for the LoTSS Survey, LOFAR as a whole and for other telescopes.

This chapter is based on the work of A.P. Mechev, T.W. Shimwell, et al. "Scalability model for the LOFAR direction independent pipeline". In: *Astronomy and Computing* 28 (2019), p. 100293. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2019.100293>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133719300290>

## 6.1 Introduction

Astronomy has entered the big data era with many projects creating petabytes of data per year. This data is often processed by complex multi-step pipelines consisting of various algorithms. Understanding the scalability of astronomical algorithms theoretically, in a controlled environment, and in production is important for making predictions for the data reduction of future projects and upcoming telescopes.

The Low Frequency Array (LOFAR) [77] is a leading European low-frequency radio telescope. The majority of LOFAR's stations are in the Netherlands, however the telescope can use stations across Europe to create ultra-high resolution radio maps. LOFAR data needs to undergo several computationally intensive processing steps before obtaining a final scientific image.

To create a broadband image, LOFAR data is first processed by a Direction Independent (DI) Calibration pipeline followed by Direction Dependent (DD) Calibration pipeline [e.g. 17–19, 83]. The goal of DI calibration is to remove effects that are constant across the target field such as radio frequency interference, contamination by bright off-axis sources and antenna gains. After this step, DD Calibration focuses on removing effects which vary across the field, such as ionospheric and beam effects. The result of these two pipelines is a science-ready image.

Our implementation of the DI LOFAR processing, `prefactor`, can be parallelized on a high throughput cluster [35]. The Direction Dependent processing, implemented in `ddf-pipeline`<sup>1</sup>, is subsequently performed on a single HPC node.

The LOFAR Surveys Key Science Project (SKSP) [24, 25] is a long running project consisting of several low frequency surveys of the northern sky. The broadest tier of the survey, LoTSS, will use more than 3000 8-hour observations to create maps with a noise levels below 100  $\mu$ Jy. We have already processed more than 500 of these observations using the `prefactor` DI pipeline [16, 40].

While the current LoTSS imaging algorithms can process data averaged by up to a factor of 64 in frequency and time, it is important to understand how LOFAR processing scales with processing parameters, such as averaging parameters. Since LOFAR data is used by multiple scientific teams, not every team can produce scientific results from data averaged by such a high factor. Users from those teams need to be able to predict the time and computational resources required to process their data, taking into account the increasing LOFAR observation rates, data sizes and scientific requirements.

<sup>1</sup>Available at <https://github.com/mhardcastle/ddf-pipeline/releases>

We study the scalability of processing LOFAR data, by setting up processing of a sample SKSP data set on an isolated node on the GINA cluster at SURF-sara, part of the Dutch national e-infrastructure [140]. We test the software performance as a function of several parameters, including averaging parameters, number of CPUs and calibration model size. Additionally, we test the performance of the underlying infrastructure, i.e. queuing and download time, for the same parameters. Finally, we compare those isolated tests with our production runs of the prefactor pipeline to measure the overhead incurred by running on a shared system.

We discover that the computationally intensive LOFAR processing steps scale linearly with data size, and calibration model size. Additionally, we find that the time taken by these steps is inversely proportional to the number of CPUs used. We discover that the time to download and extract data on the GINA cluster is linear with size up to 32GB, but becomes slower beyond this data size. We also find that the queuing time on the GINA cluster grows exponentially for jobs requesting more than 8 CPUs. We validate these isolated tests with production runs of LOFAR data from the past six months. We combine all these tests into a single model and show its prediction power by testing the processing time for different combinations of parameters. Finally, we discuss the utility of our method, the results in this work and applications to the SKSP projects, the broader impact of our results to LOFAR processing and the applications for other large astronomical surveys. The major contributions of this work can be summarized as:

- A model of processing time for the LOFAR Direction Independent Calibration Pipeline.
- A model of queuing time and file transfer time which is used by current or future jobs processed on the GINA cluster.
- Quantification of overheads incurred when processing in production.
- Validation of our methods with discussion of future applications.

We introduce LOFAR processing and other related work in Section 6.2 and describe our software setup and data processing methods in Section 6.3. We present our results and performance model in Section 6.4 and discussions and conclusions in Section 6.5.

## 6.2 Related Work

In previous work, we have parallelized the Direction Independent LOFAR pipeline on a High Throughput infrastructure [35]. While this parallelization has helped accelerate data processing for the SKSP project, creating a performance model of our software is required if we are to predict the resources taken by future jobs. This model will be particularly useful in understanding how processing parameters will affect run time.

Performance modelling on a distributed system is an important field of study related to Grid computing. A good model of the performance of tasks in distributed workflows can help more efficiently schedule these jobs on a Grid environment [141]. The performance modelling systems require knowledge of the source code and an analytical model of the slowest parts of the code [142]. Many systems exist to model the performance of distributed jobs [142–145], with some employing Black Box testing [146, 147] or tests on scientific benchmark cases [148]. Such performance analysis does not require intimate knowledge of the software and can be applied on data obtained from processing on a Grid infrastructure.

Empirical modelling is useful in finding performance bugs in parallel code [149] and modelling the performance of big data architectures [150]. The insights from these models are used to optimise the architecture of the software system or determine bottlenecks in processing. Here, we use empirical modelling to determine how the LOFAR `prefactor` performance scales with different parameters.

## 6.3 Processing Setup

Using the LOFAR software installation described in [35], we processed a typical LOFAR SKSP observation<sup>2</sup>, while changing the averaging rate in time and frequency. Changing these averaging parameters will change the final data size (with the data sizes studied shown in Table 6.1). We test the processing time for different averaging parameters by running 15 runs per parameter step.

The data used by the LOFAR surveys is archived at a time resolution of 1 second intervals and frequency resolution of 16 channels per Subband (equivalent to 12kHz channel width). While some of the processing steps such as flagging of

<sup>2</sup>LOFAR Observation ID L658492, co-ordinates [17h42m21.785, +037d41m46.805] observed by the LOFAR High Band Array for 8 hours between 2018-06-20 and 2018-06-21.

Radio Frequency Interference and removal of bright off-axis sources produce better results when performed on the high-resolution data, later steps can be performed on averaged data with little impact on the final product quality. To speed up processing, the raw data is averaged in time and frequency, decreasing the input data size to later tasks. The main aims of the LoTSS survey can be accomplished if the final data products from the `prefactor` pipeline are averaged to a time resolution of 8 seconds per sample and frequency of 2 channels per subband. These averaging parameters correspond to a reduction in data size by a factor of 64. Nevertheless, other science cases require less averaging of the data. Our aim is to understand how the processing of this larger data will scale. To measure the scalability of processing, we measure the performance of the `prefactor` pipeline for data sizes between the raw data of 64GB/subband and the averaged data of 1GB/subband. The tested data sizes and parameters are shown in table 6.1, and discussed in Section 6.4.1.

We performed the scalability tests on a dedicated node of the SURFsara GINA cluster, f18-01. The node is a typical hardware node used by our production LoTSS processing, however it is dedicated for the tests in order to ensure there is no contamination by other software. The node is described in Section 6.3.3.

We processed the sample data set with the LOFAR `prefactor` pipeline. The `prefactor` version used was the same as we use for the LOFAR SKSP broadband surveys [40]. This software consists of several steps executed in sequence, shown graphically in Figure 6.1. The important `prefactor` steps are as follows. The `predict_ateam` and `ateamcliptar` steps predict the contamination by bright off-axis sources and remove these effects respectively. The `dpppconcat` step is responsible for concatenating 10 Subband into a single file which is in turn calibrated. The step `gsmcal_solve` is responsible for calibration of the data against a model of the radio sky. The solutions produced by `gsmcal_solve` are used by `gsmcal_apply` and applied to the scientific observation.

### 6.3.1 Processing Metrics

The goal for our scalability model is to understand the effect of several parameters on the job completion time of LOFAR software. We do this by testing the processing time for various values of data size, number of CPUs used and sky model size.

The slowest step of the `prefactor` pipeline is the `gsmcal_solve` step, which performs the gain calibration against a model of the radio sky. This step operates on a concatenated data set that consists of 10 Subbands. We obtain the

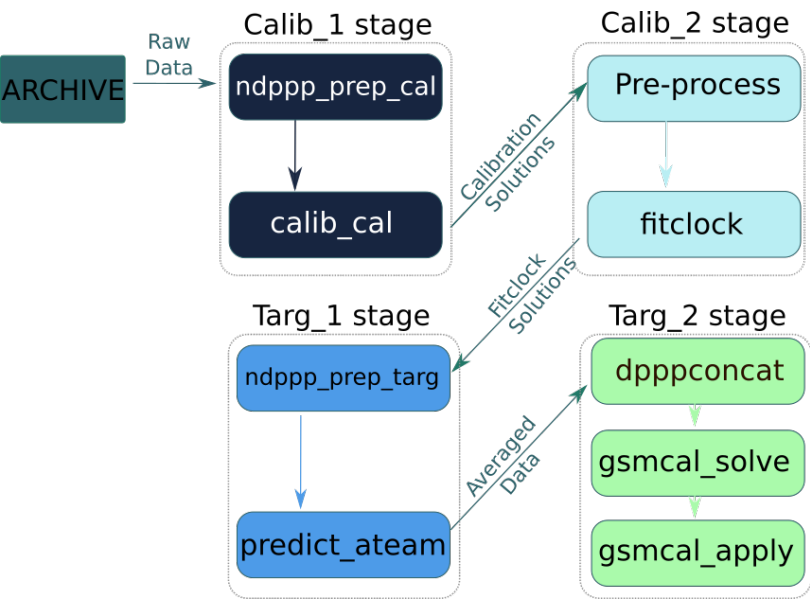


Figure 6.1: The major steps of the prefactor DI pipeline.

| Averaging ratio | Time averaging parameter (sec) | Channels per Subband | Averaged Size (Gb) |
|-----------------|--------------------------------|----------------------|--------------------|
| 64x             | 8                              | 2                    | 1.235              |
| 32x             | 4                              | 2                    | 2.459              |
| 16x             | 2                              | 2                    | 4.906              |
| 8x              | 1                              | 2                    | 9.802              |
| 4x              | 1                              | 4                    | 18.00              |
| 2x              | 1                              | 8                    | 36.72              |
| 1x              | 1                              | 16                   | 66.88              |

Table 6.1: Averaging parameters and final data sizes tested for the sample LOFAR SKSP observation. The raw data is 64 GB per subband. The LOFAR SKSP data processing uses averaging parameters of 8 seconds and 2 channels per subband. This reduces the raw data by a factor of 64. We highlight the data size used in the LOFAR SKSP survey.

calibration model through the TGSS sky model creator<sup>3</sup>. By default, this service creates a text file describing the sky-model from the TGSS survey [151] as a combination of Gaussians and point sources. By default, it sets a threshold of sources brighter than 0.3 Jy. Lowering this threshold creates longer sky-model files with more faint sources, while increasing it will return only the few brightest sources.

<sup>3</sup>Accessible at the TGSS ADR portal.

Since sky model calibration requires converting the sky-model into modelled visibilities[e.g. 99, 152, 153], a longer sky model will increase the time taken to gain calibrate a data set. We created seven sky models with a flux cutoff ranging between 0.05 Jy and 1.5 Jy. The number of sources in the resulting models are listed in Table 6.2. For production<sup>4</sup>, we used the minimum sensitivity parameters for model 3.

It is important to note that the complexity and accuracy of the sky model depend on the direction of observation and the ionospheric conditions in which the observation was performed. As such, our test is only a heuristic for predicting the run-time based on the calibration model length. Additionally, it is notable that the number of sources is exponentially dependent on the minimum sky model sensitivity (seen in Figure 6.2, more in [17, 151]). According to this relationship, even a modest decrease in sensitivity cutoff can significantly decrease the size of the model.

| Sky model # | min sensitivity | # sources |
|-------------|-----------------|-----------|
| model 1     | 0.05 Jy         | 809       |
| model 2     | 0.1 Jy          | 503       |
| model 3     | 0.3 Jy          | 180       |
| model 4     | 0.5 Jy          | 96        |
| model 5     | 0.8 Jy          | 49        |
| model 6     | 1.0 Jy          | 34        |
| model 7     | 1.5 Jy          | 16        |

Table 6.2: List of test sky models. Model 3 is created with the parameters used in our production processing of LOFAR data. All models include objects within 5 degrees from the centre of the pointing.

Finally, the number of CPU cores (henceforth just ‘CPUs’) used by each step is a parameter that can be optimised for the entire pipeline. While increasing the number of CPUs can make some steps run faster, requesting jobs that reserve a large number of CPUs will take longer to launch on shared infrastructure. In order to understand the interplay between these effects, we study the queuing time and processing time as a function of the number of CPUs. For the parameter steps we choose to test 1, 2, 3, 4, 8 and 16 CPUs.

### 6.3.2 Infrastructure Performance

In production, we run hundreds of LoTSS jobs on a cluster supporting several different use cases. The requested resources on this cluster are allocated by a job queue,

<sup>4</sup>The query used to obtain model 3 is at the following link [http://bit.ly/tgss\\_model](http://bit.ly/tgss_model)

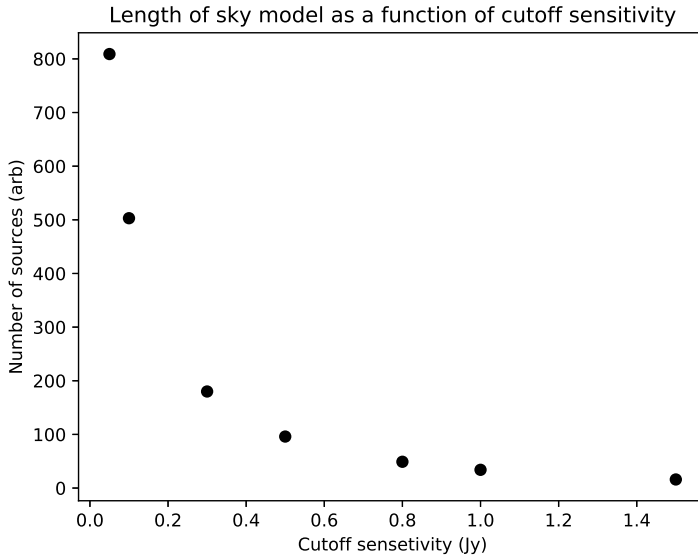


Figure 6.2: The size of the sky model (measured as the number of sources) increases exponentially as we decrease the flux cutoff of the model (i.e. increase the sensitivity).

in our case implemented by the glite workload management system [154]. As queuing jobs can take a significant amount of time, we test the queuing time as a function of the number of requested CPUs. In order to do that, we create test jobs that log the launch time and submit them, requesting 1, 2, 3, 4, 8 and 16 CPUs. We run 10 to 15 tests for each parameter step to ensure that we capture system variability at different times of day during the week and the weekend.

Besides queuing, time is also spent during downloading and unpacking data, as well as packing and uploading the results. Despite using no compression to pack the data, untarring and tarring large files still takes time depending on the system workload. We measure the time taken to transfer and unpack data of different sizes. The data sizes we chose were 0.5GB, 1GB, 2GB, 4GB, 8GB, 16GB, 32GB and 64GB. As our largest data sets are 64GB and our smallest results are  $\sim 0.2$ GB, these values span a realistic range expected for LOFAR data processing. We test this by uploading mock data to the dCache storage pool at SURFsara and launching a small 1 CPU job on the cluster, which downloads and untars the data and logs the start time of each step. We present the results of this test in the next section.

## Software Versions

For the current test, we use the LOFAR software stack, version 2.20.2 [95]. This software was compiled on a virtual machine and distributed using the CERN CVMFS virtually mounted file system [58]. We use this software version and distribution method as it is the same software version and distribution used to process the data for the LOTSS Data Release 1.

### 6.3.3 Test Hardware

We test the LOFAR software on a reserved node on the SURFsara GINA cluster. The node, f18-01 has 348 GB of RAM, 3TB of scratch space<sup>5</sup>. The CPU is an Intel Xeon(R) Gold 6148 CPU with 40 cores clocked at 2.40GHz. As this hardware node was reserved, there was no other scientific processing aside from our tests, meaning there was no resource contention aside for that inherent in the LOFAR software. In the results section, we compare these isolated runs with processing results over the past two years.

## 6.4 Results

Using a test data set, we tested the LOFAR `prefactor` target pipeline on the SURFsara GINA cluster. First we will present the tests done in an isolated environment and then compare them to the run time in production on a shared infrastructure. We will integrate all the results in a complete model which can be used to predict processing time for a variety of parameters. Finally, we will make some predictions on the run time of our processing based on the model and validate these predictions.

Since we are processing a sample data set in the context of the LOFAR Surveys project, we will compare these tests with the production runs of our pipeline. In production, we run the `gsmcal_solve` step with a data size of 1GB, a sky model with 180 sources and 8 CPUs.

### 6.4.1 Isolated Environment tests

We first tested the LOFAR software in isolation in order to determine the scalability of processing time in terms of data size. We run the entire `prefactor` target pipeline

---

<sup>5</sup>More detailed specifications are at the GINA specification page linked here

$$T_{predict\_ateam} = 5.19 \times 10^{-8} \mathcal{S} + 4.20 \times 10^1 \quad (6.1a)$$

$$T_{ateamcliptar} = 4.57 \times 10^{-9} \mathcal{S} - 8.42 \times 10^0 \quad (6.1b)$$

$$T_{dpppconcat} = 3.51 \times 10^{-8} \mathcal{S} + 4.20 \times 10^1 \quad (6.1c)$$

$$T_{gsmcal\_solve} = \begin{cases} 7.38 \times 10^{-7} \mathcal{S} - 8.20 \times 10^1 & |\mathcal{S} \leq 1.6 \times 10^{10} \\ 1.04 \times 10^{-6} \mathcal{S} - 4.04 \times 10^3 & |\mathcal{S} > 1.6 \times 10^{10} \end{cases} \quad (6.1d)$$

$$T_{gsmcal\_apply} = 2.07 \times 10^{-8} \mathcal{S} - 1.38 \times 10^1 \quad (6.1e)$$

Equation 6.1: Equations describing the processing time of five **prefactor** steps as a function of the input data size ( $\mathcal{S}$ ) in bytes.

which removes Direction Independent Calibration errors from a LOFAR science target. In the following sections, we present the models obtained from these tests.

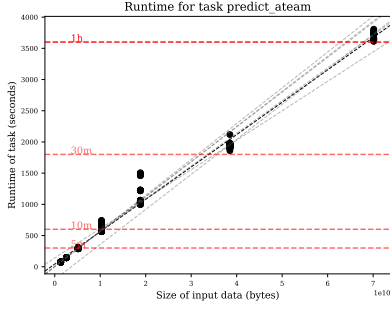
### Input Data Size

LOFAR data can be averaged to different sizes based on the scientific requirements. Smaller data sets are processed faster, so it is important to understand the effect of data size on processing time as measured by wall-clock time. We show the processing time for our test data set, averaged to different sizes for several **prefactor** steps in Figures 6.3a- 6.3d and 6.4. We run this test using 8 CPUs. The figures also show linear fits for consecutive pairs of parameter steps, in gray dashed lines, used to help guide the selection of the parametric model.

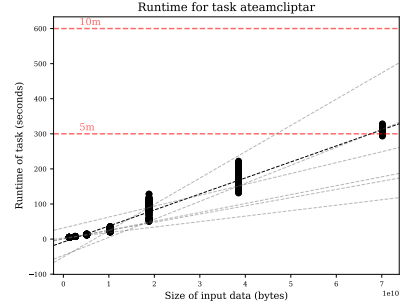
All of the steps show a linear behavior with respect to input data size, while the `gsmcal_solve` step is best fit by two linear relationships, for data smaller and larger than 16 GB. The linear fit to the run times are shown in Equations 6.1a-6.1e. The equations show the processing time as a function of the data size ( $\mathcal{S}$ ), with the slope in the units of seconds/byte. The fits are also shown in Figures 6.3a to 6.4 as a black dashed line.

### Calibration Model Size

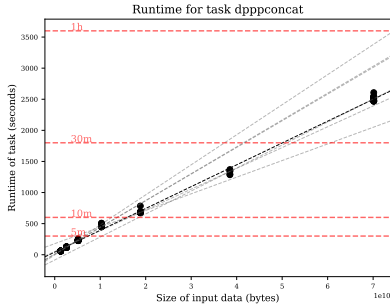
To test the effect of the calibration model size on run time, we test our calibration with several different lengths of the sky model file. We create these models by changing the minimum sensitivity using values ranging from 0.05 Jy to 1.5 Jy. The



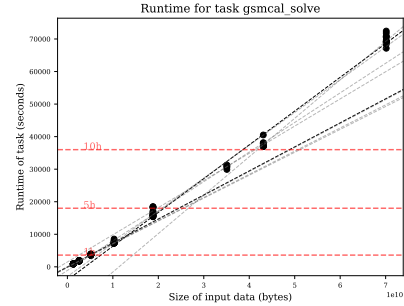
(a) Tests of the `predict_ateam` step for input data size ranging from 1GB to 64 GB. This step calculates the contamination from bright off-axis sources. Dashed lines are shown connecting each pair of points, to highlight the trend. We can see that the data can be described well by a linear model. We show the model in Equation 6.1a in black.



(b) Tests of the `ateamcliptar` step for input data size ranging from 1GB to 64 GB. This step removes the contamination from bright off-axis sources. We can see that the data fits a linear model. We show the model in Equation 6.1b in black.



(c) Tests of the `dpppconcat` step for input data size ranging from 1GB to 64 GB. This step concatenates 10 files into a single measurement set. We can see that the data fits a linear model. We show the model in Equation 6.1c in black.



(d) Tests of the `gsmcal_solve` step for input data size ranging from 1GB to 64 GB. This step performs gain calibration of the concatenated data set against a sky model. It is the slowest and most computationally expensive `prefactor` step. We fit two linear models, for data below 16GB and above 16GB. We can see the two models, shown in (Equation 6.1d) as two black dashed lines, intersecting at 1GB.

Figure 6.3: Plots of the run time as a function of input data size

most sensitive model (0.05 Jy) had 809 sources, while the 1.5 Jy model had only 16 sources.

Figure 6.5 shows that the calibration time is directly proportional to the

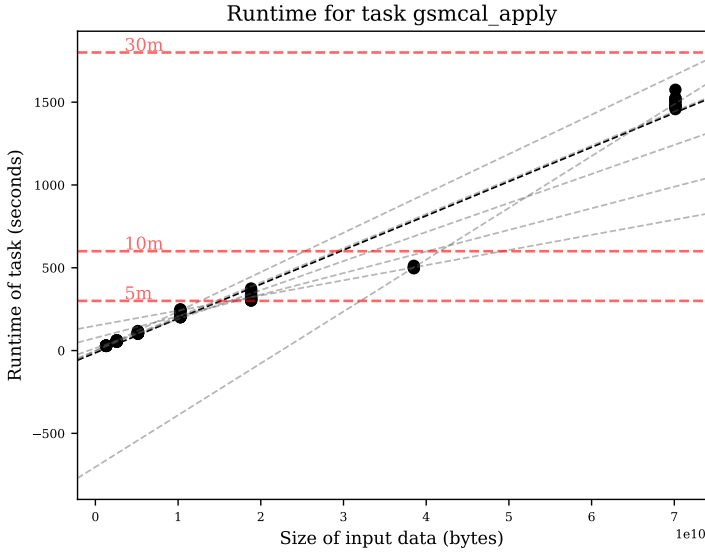


Figure 6.4: Tests of the `gsmcal_apply` step for input data size ranging from 1GB to 64 GB. This step applies the calibration solutions to the data. We can see that the data fits a linear model, described in Equation 6.1e, as the black dashed line.

length of the sky model. Figure 6.6 shows the run time as a function of the processing parameter: the cutoff sensitivity. As the relationship between the number of sources and cutoff sensitivity is a power law, here we see the same relationship holding for processing time.

6

We model the processing time as a function of the cutoff frequency using a power law, and fit the data to the function  $y = \alpha \cdot \mathcal{F}^{-k}$ . Our fit found the best model to be shown in Equation 6.2, where  $\mathcal{F}$  value is the cutoff flux in Jansky and  $T$  is the run time in seconds.

We show four images made from data sets in Figure 6.7. The top left image is calibrated with a 0.05Jy and the other three show the difference between the top left image and the images created from the 0.3Jy, 0.8 Jy and 1.5 Jy data. The statistics for the four images, taken from the regions in green on Figure 6.7) are shown in Table 6.3. We discuss the implication and caveats of these results in Section 6.5.

<sup>6</sup>Using the command `wsclean -absmem 50 -niter 3 -size 4096 4096 -scale 20asec`

$$T = 1185 \cdot \mathcal{F}^{-0.854} \quad (6.2)$$

Equation 6.2: Processing time for the `gsmcal_solve` step as a function of the flux cutoff of the calibration model ( $\mathcal{F}$ ) in Jansky.

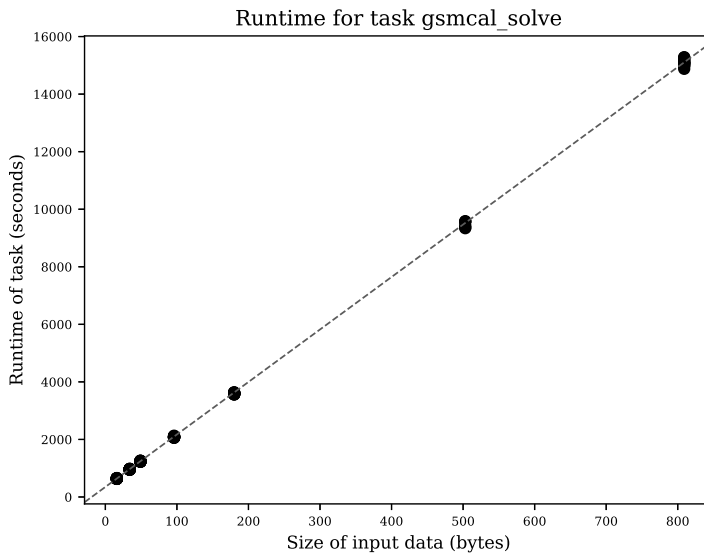


Figure 6.5: The processing time of the `gsmcal_solve` step is linear with the size of the sky model as measured by the number of sources.

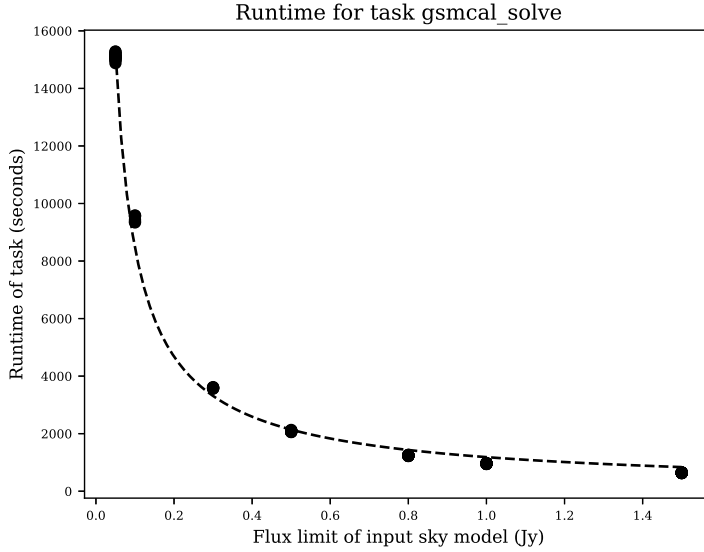


Figure 6.6: The run time of the `gsmcal_solve` step as a function of the cutoff sensitivity is not linear. As shown in Figure 6.2, the number of sources increases exponentially as the minimum sensitivity decreases. The dashed line shows the model fitted in Equation 6.2.

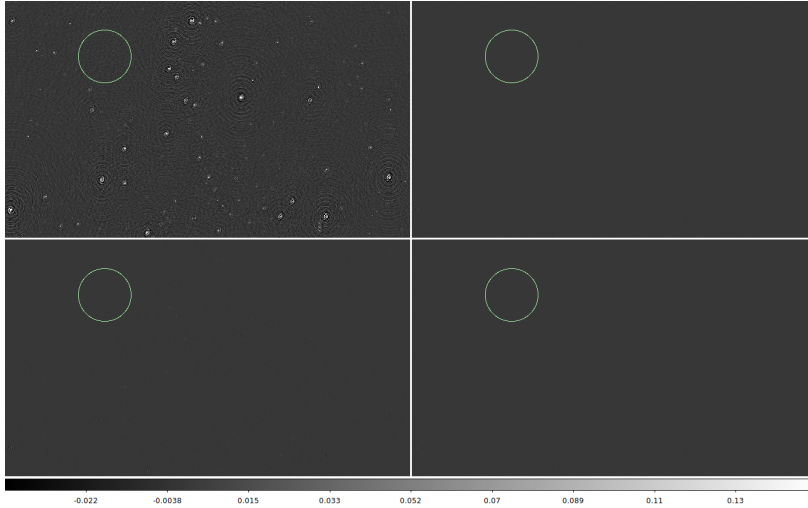


Figure 6.7: Four images made using the `wsclean` software [132] from the data set<sup>6</sup>. The four images were calibrated with sky models of various flux cutoffs ranging from 0.05Jy (top left) to 1.5Jy (bottom right). Flux statistics for the green regions in the four images are listed in Table 6.3. The top right and bottom two quadrants show the pixel difference between the 0.05Jy image and the 0.3Jy, 0.8Jy and 1.5Jy images respectively. The four images are all on the same scale. The green region shows the same area in all four quadrants.

| Calibration Model<br>Flux Cutoff | # of sources | RMS Noise (Jy) |
|----------------------------------|--------------|----------------|
| 0.05Jy                           | 809          | 0.00402834     |
| 0.3 Jy                           | 180          | 0.00402311     |
| 0.8 Jy                           | 49           | 0.00404181     |
| 1.5 Jy                           | 16           | 0.00410204     |

Table 6.3: Statistics for an empty region for the four images shown in Figure 6.7. The 0.3Jy model, here shown shaded in gray, is the one used in production.

$$T = 503.37 + \frac{3062.6}{\mathcal{N}} \quad (6.3)$$

Equation 6.3: Processing time for the `gsmcal_solve` step as a function of ( $\mathcal{N}$ ), the Number of CPUs used by the process.

### Number of CPUs

One further parameter that can be optimised is the number of CPUs requested when the job is launched. We investigated the processing speedup as a function of the number of CPUs for the `prefactor` target pipeline. From the steps tested, only the `gsmcal_solve` step shows a significant speedup as the number of CPUs is increased. The run time of this step is an inverse power law with respect to the number of CPUs as seen in Figure 6.8. Unlike the solving step, the step applying the calibration solutions (`gsmcal_apply`) is constant in time with respect to the number of CPUs as seen in Figure 6.9. The difference in performance is most likely because the `gsmcal_apply` code uses a parallel for loop to calculate antenna gains while `gsmcal_solve` does not.

We fit a model with processing time inversely proportional to the number of CPUs used. We show the resulting model in Equation 6.3, with the parameter ( $\mathcal{N}$ ) being the number of CPUs used.

### 6.4.2 Queueing Tests

Aside from the performance of the LOFAR software, we measured the queueing time at the GINA cluster, as a function of the number of CPUs requested. This data was obtained between 16 Nov 2018 and 10 Dec 2018 for 1, 2, 3, 4, 8, and 16 CPUs per job. A histogram of the queueing time for these jobs is shown in Figure 6.10. Statistics for these runs are in Table 6.4. We use the 75th percentile of the queueing

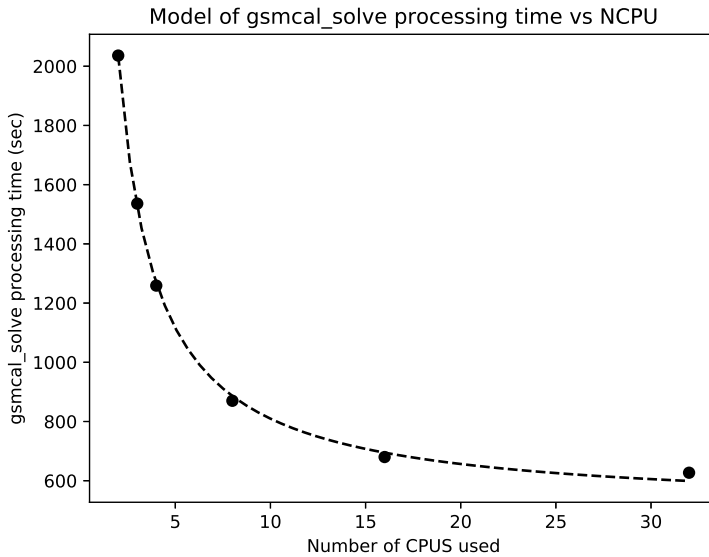


Figure 6.8: The processing time of the `gsmcal_solve` step decreases exponentially with the number of CPUs requested. The model in Equation 6.3 is shown in a dashed line. As this is a  $1/x$  model, it shows diminishing returns past 16 CPUs.

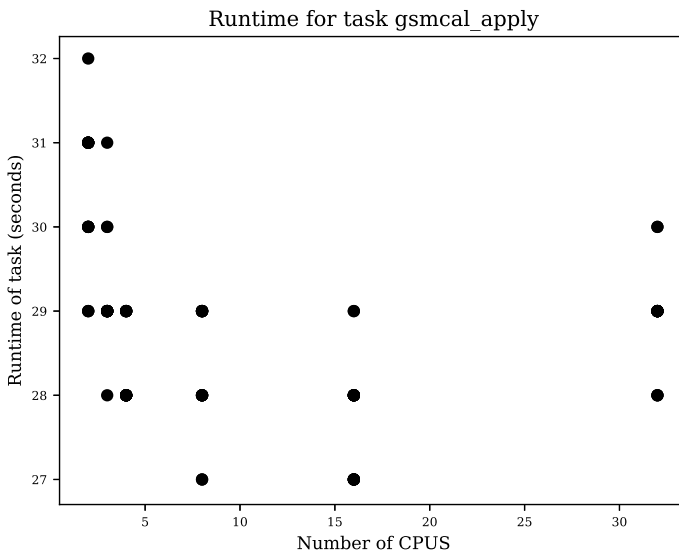


Figure 6.9: The step that applies the calibration solutions, `gsmcal_apply`, does not show a speedup when run on multiple cores, as all runs take roughly 30 seconds to complete.

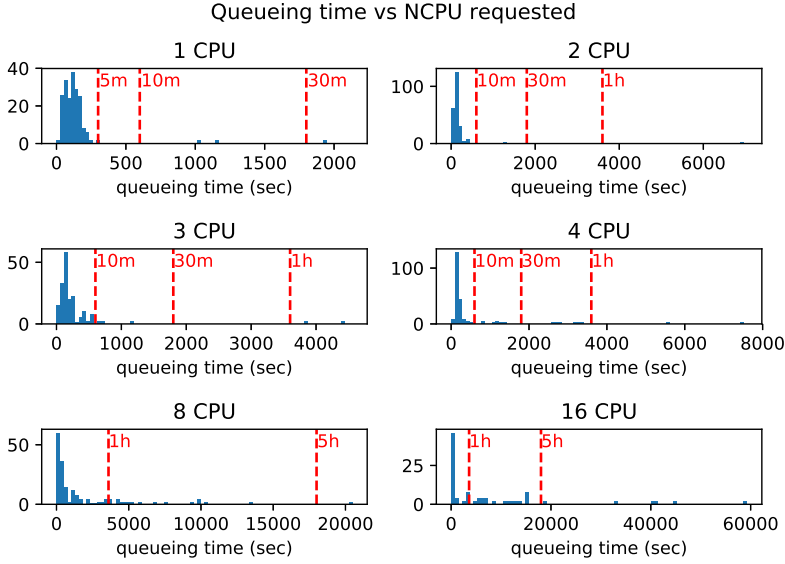


Figure 6.10: Test randomly submitting jobs to the GINA with different number of requested CPUs. The long tail for 8 and 16 CPU jobs shows that some jobs can take several hours to launch.

$$T = \begin{cases} 49.3 \cdot \mathcal{N} + 120 & |\mathcal{N}| \leq 4 \\ 726 \cdot \mathcal{N} - 3071 & |\mathcal{N}| > 4 \end{cases} \quad (6.4)$$

Equation 6.4: The model for the queueing time as described by two linear models.

time for each parameter step to fit a model. This scenario will include 75% of runs and is a good trade-off between ignoring and including outliers.

We fit two linear models for this queueing time. One model for 1-4 CPUs and one for 4-16 CPUs. The model, as a function of the number of CPUs  $\mathcal{N}$  is in equation 6.4. The two models are plotted against the 75th percentile of the queueing times (last column in Table 6.4) in Figure 6.11.

### 6.4.3 Transfer and Unpacking Time

We tested the downloading and unpacking time for data sizes ranging from 512MB to 64GB. We discovered that the unpacking of files below 64GB scaled linearly with file size, however unpacking individual data sets larger than 16GB becomes considerably slower than downloading it.

| NCPU requested | Mean time (sec) | Median time (sec) | 75th percentile (sec) |
|----------------|-----------------|-------------------|-----------------------|
| 1 CPU          | 150.5           | 116.2             | 154.1                 |
| 2 CPU          | 201.1           | 125.8             | 165.8                 |
| 3 CPU          | 296.2           | 152.0             | 243.0                 |
| 4 CPU          | 498.9           | 167.7             | 233.7                 |
| 8 CPU          | 1944.2          | 428.4             | 2142.4                |
| 16 CPU         | 7079.0          | 696.4             | 8750.6                |

Table 6.4: Statistics for queuing time for different values of CPUs requested. Queueing time for jobs requesting less than 8 CPUs is typically less than five minutes. It can drastically increase for larger jobs.

$$T = 5.918 \times 10^{20} \cdot \mathcal{S}^{2.336} \quad (6.5)$$

Equation 6.5: Model of the downloading and extracting time as a function of the data size ( $\mathcal{S}$ ) in bytes.

Figure 6.12 shows the histogram of the download tests, and Figure 6.13 displays the tests as a function of data size. Both figures show that extracting of the 32 and 64GB data sets has more slow outliers than the downloading of this data.

We fit a power law model to the time taken to transfer and unpack the data. In this case, we also consider the 75th percentile of these times in order to capture the majority of runs and ignore outliers. The plot of the data and our model can be seen in Figure 6.13 and the model is in Equation 6.5, as a function of the input data size,  $\mathcal{S}$  in gigabytes.

## 6

### 6.4.4 Comparison with production runs

Over the past two years, the LOFAR software has been running in production and collecting data on run time for each processing step. We have saved detailed logs for these runs starting in July 2018. We can compare this to the isolated model in order to determine the overhead incurred by processing LOFAR data on shared nodes.

Using the logs recorded by our processing launcher <sup>7</sup>, we made plots showing the processing time for the downloading and extracting, and for the slowest steps, `ndppp_prepcal` and `gsmcal_solve`. The results are shown in Figures 6.15 and 6.16. We include predicted extract times from Section 6.4.3 as vertical dashed lines for both plots. The agreement between our model and production runs are an encouraging result for future software performance modelling.

<sup>7</sup>GRID\_PiCaS\_Launcher, [https://github.com/apmechev/GRID\\_picastools](https://github.com/apmechev/GRID_picastools)

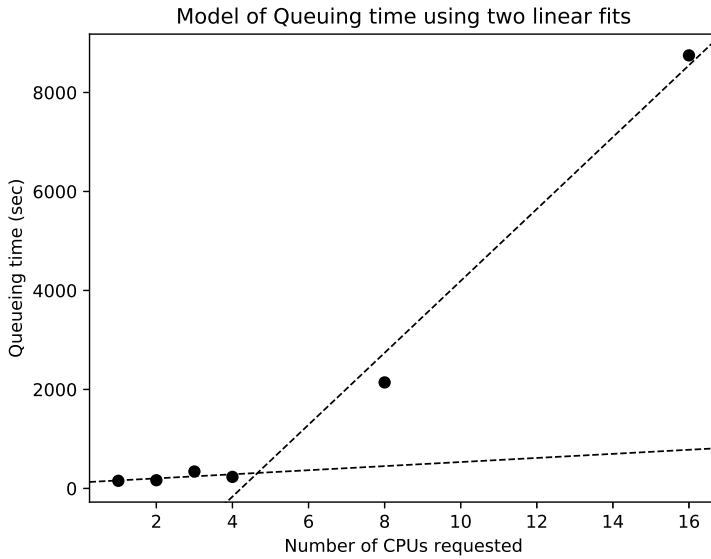


Figure 6.11: The queuing model built from two linear fits to the queuing times. We use the 75th percentile of the queuing data as an upper bound of job queuing.

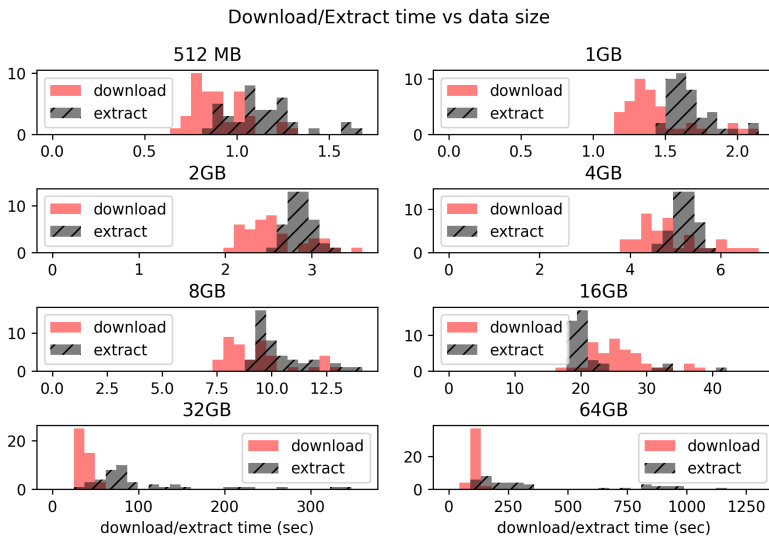


Figure 6.12: A histogram of the download and extracting times of multiple data sizes on the GINA worker nodes. Download and extract times are comparable for data up to 8GB, however above that, the extracting time dominates.

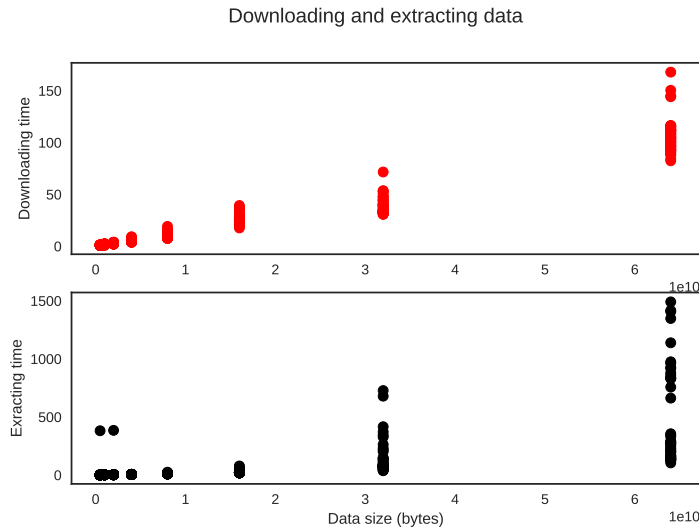


Figure 6.13: A scatter plot of the download and extracting times of multiple data sizes on the GINA worker nodes. The difference between download and extract time for the 32 and 64 GB data sets can be seen.

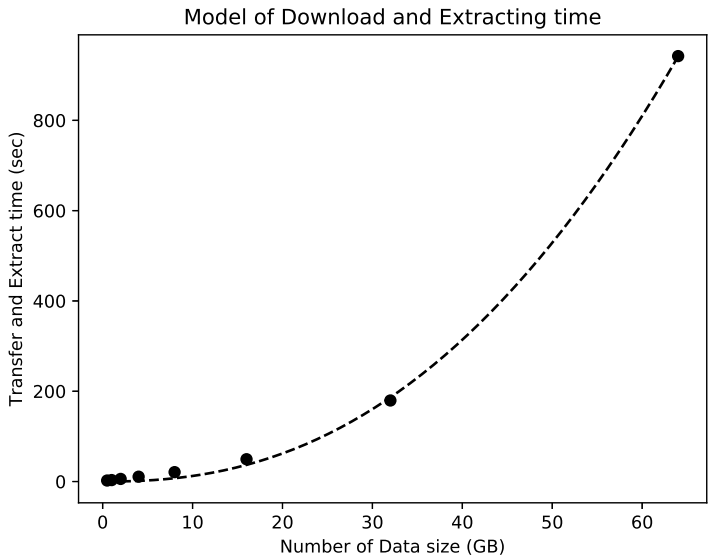


Figure 6.14: Fit of an exponential model to the Download and Extraction time for different data sizes. For the transfer overhead, we took the 75th percentile from the data shown in Figure 6.12. The model in Equation 6.5 is shown in a dashed line.

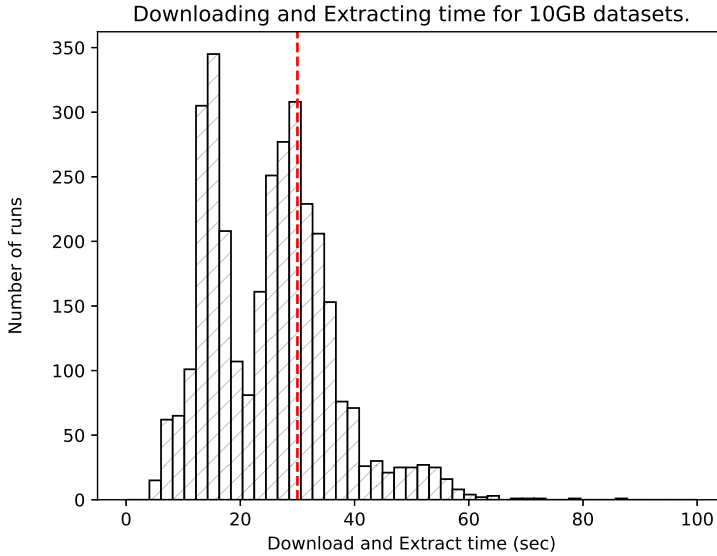


Figure 6.15: Downloading and extracting time for 10 1GB data sets performed in our production environment. Data from this test ranges from July 2018 to January 2019. The dashed red line shows the prediction obtained from section 6.4.3. We see a bimodal distribution corresponding to 10 GB data (right peak) and data averaged further to 5GB (left peak).

Finally, we present Figure 6.17 which shows a comparison of `gsmcal_solve` run times and our model's prediction for a 1GB data set. Figure 6.18 plots the processing time vs data size for these production runs and includes the model from Equation 6.1d. The significant overhead incurred on a shared infrastructure can be noted.

### 6.4.5 Complete Scalability Model

To incorporate all our data into a complete model, we consider the slowdown of each parameter as a multiplier to the time taken to process our base run. We incorporate the models for each parameter above for the model of the run time. We add the transfer and queuing time to the processing time to obtain a final function of all our parameters. We can use this function to predict the processing time for an arbitrary data set.

The final performance model for the slowest steps, `gsmcal_solve`, `dpp-pconcat` and `predict_ateam` are in Equation 6.6.

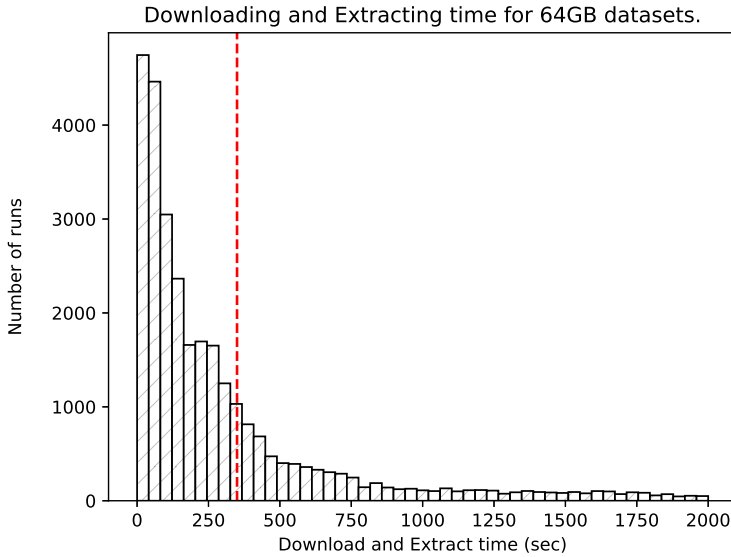


Figure 6.16: Downloading and extracting time for a 64GB data set performed in our production environment. Data from this test ranges from 07/2018-01/2019. The dashed red line shows the prediction obtained from Figure 6.3d in Section 6.4.3.

## 6.5 Discussions and Conclusions

The goal of this work is to understand the performance of the LOFAR Direction Independent Pipeline as processing parameters are changed. We modify several parameters and compare the wall clock time taken to process the data. Finally, we study data from queuing jobs and downloading data in order to fully model infrastructure overheads. We compare our model with past runs of the software and discuss the results and implications. We discuss the utility of this model for current and upcoming LOFAR projects. Finally, we note the effectiveness of this modelling technique for understanding the performance of large-scale processing of astronomical data.

### 6.5.1 Software Performance

We performed several tests to determine the scalability of LOFAR processing with respect to several parameters. We outline our findings below as well as their implications for processing in the context of the LOFAR surveys and other LOFAR projects.

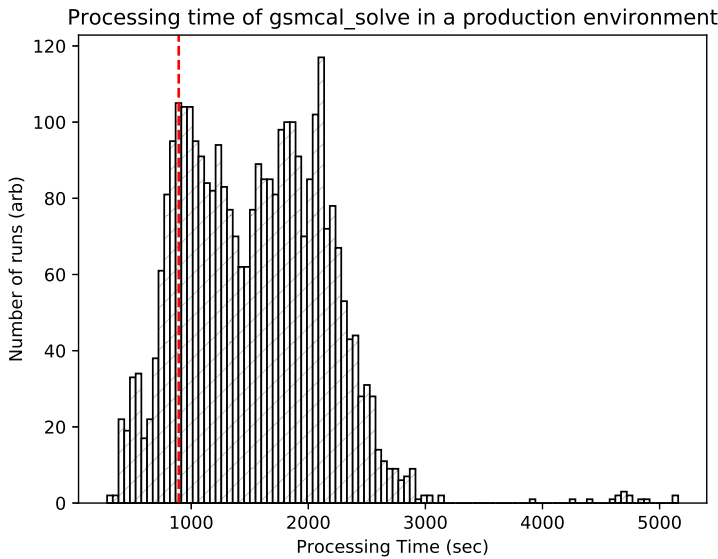


Figure 6.17: Processing time for the `gsmcal_solve` step in a production environment. Data from this test ranges from 07/2018-01/2019. The dashed red line shows the prediction for a 1GB run, obtained from section 6.4.3. We see two distributions, which correspond to data averaged to 1GB and 512 MB. It should be noted that the left peak corresponds to 512MB data, as seen in Figure 6.18.

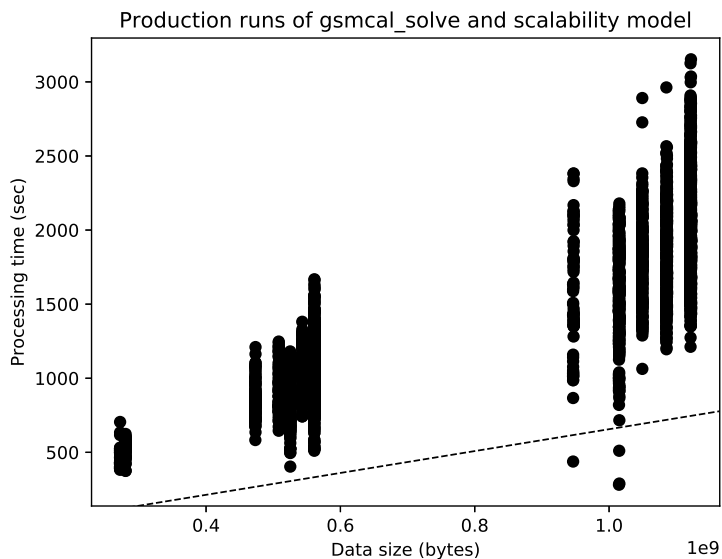


Figure 6.18: The scalability model for processing data through the `gsmcal_solve` step, shown in a dashed line. The scatter plot shows the performance for production runs of this step between July 2018 and January 2019. The two large clusters are for data products that are 1.0 and 2.0 GB respectively.

$$t_{infrastructure} = \begin{cases} 49.3 \cdot \mathcal{N} + 120 & |\mathcal{N} \leq 4 \\ 726 \cdot \mathcal{N} - 3071 & |\mathcal{N} > 4 \end{cases} + 2 \cdot 0.056 \cdot \mathcal{S}^{2.336} \quad (6.6a)$$

$$t_{gsmcal\_solve} = t_{infrastructure} + [3566 \cdot \frac{1}{3.012} \mathcal{F}^{-0.854} \cdot (0.1412 + \frac{0.8589}{\mathcal{N}})] \cdot \begin{cases} 7.38 \cdot 10^{-7} \mathcal{S} & \mathcal{S} \leq 16 \\ 1.04 \cdot 10^{-6} \mathcal{S} & \mathcal{S} > 16 \end{cases} \quad (6.6b)$$

$$t_{dppconcat} = t_{infrastructure} + 3.51 \times 10^{-8} \mathcal{S} + 4.20 \times 10^1 \quad (6.6c)$$

$$t_{predict\_ateam} = t_{infrastructure} + 5.19 \times 10^{-8} \mathcal{S} + 4.20 \times 10^1 \quad (6.6d)$$

Equation 6.6: Model of the total time of the most computationally expensive steps for the parameters  $\mathcal{N}$ , Number of CPUs;  $\mathcal{S}$ , Size of data in bytes and  $\mathcal{F}$ , cutoff calibration model flux in Jansky. These models include processing times, as well as infrastructure overheads. As the model for the queuing, downloading and uploading time does not change for different processing steps, we decide to keep it separate for clarity. The complete scalability models for the rest of the steps can be derived similarly, however are omitted here as they consist of the minority of processing time for LOFAR DI processing.

## Data Size

We tested broadband LOFAR data ranging in size by a factor of 64, and discover that all our processing steps scale linearly in time with respect of the input data size. We learn that for input data above 16GB, the slope of our scaling relation is higher than for the smaller data sets. The linear scaling of our processing suggests that projects interested in processing massive LOFAR data sets can scale well in terms of processing time.

As the calibration step concatenates 10 input Subbands, data larger than 16GB shows a higher slope (Figure 6.3d), meaning they take longer to process than smaller data. This can be attributed to the large memory requirement for data larger than 160GB which is likely due to the additional memory requirements of the minimization algorithm. Splitting the performance model in two also helps make a

more accurate processing time prediction as fitting a single linear model would have a large negative y-intercept, predicting negative processing times for data smaller than 2GB.

Our analysis shows the following results. Overall, the slowest step was the `gsmcal_solve` step, and its run time scales more strongly with data size than the other steps (equation 6.1d has the steepest slope). This suggests that as data sizes increase, `gsmcal_solve` will increasingly dominate the processing time over the other steps (As seen in Figures 6.3a and 6.3d). This effect is especially prominent for data larger than 16GB (160GB when 10 Subbands are concatenated). As such, it is recommended to avoid calibration of data larger than 160GB. This limitation suggests that science requiring operations on non-averaged data sets, such as long-baseline imaging, will require significant computational requirements for high-fidelity images.

### Calibration Model Size

We discover that the calibration time scales linearly in as a function of the length of the calibration model, however as a power law with respect to the model's cutoff sensitivity. This is because of the (expected) power law relation between the number of sources and cutoff sensitivity, seen in Figure 6.2. We can use this discovery to accelerate the processing time by increasing the flux cutoff to the LOFAR direction independent calibration to 0.5 Jy. Doing so will execute the calibration step in 60% of the time, saving 83 CPU-h per run. Over the remaining 2000 prefactor runs left in the LOTSS project, this change in sensitivity will save more than 167k CPU-hours.

Figures 6.7 show a data set calibrated with sky models with cutoff sensitivities listed in Table 6.3, and figures 6.19 and 6.21 show the calibration solutions obtained by calibrating with sky models of cutoff ranging from 0.05 Jy to 1.5 Jy. These results suggest that performing gain calibration with less complex, and thus smaller, calibration models will not degrade image and solution quality while taking less than 20% of processing time. Table 6.3 also confirms this result.

While this result is encouraging, there are caveats suggesting future study is required. The results we present are for a single observation and has not been processed through the Direction Dependent Calibration pipeline. This pipeline produces final images used for scientific research. Future work will need to confirm that smaller calibration models used in the `prefactor` pipeline do not degrade the

quality of these final images. Nevertheless, if this result holds, the calibration model threshold can be decreased for a wide range of LOFAR projects, significantly saving processing time and computing resources.

### Comparison with production runs

When comparing our model's prediction with real processing runs over the past six months, we note that there are considerable overheads when running on a shared infrastructure vs. when processing data on an isolated (Figure 6.18). The overhead in processing is roughly a factor of two-three from our model. This discrepancy suggests that a model for `gsmcal_solve` needs to be built using data when running on a shared environment, to better predict processing time.

### 6.5.2 Infrastructure Performance

We tested downloading and extracting LOFAR data of various sizes. Both downloading and extracting are shown to be linear in time with respect to the data size for data up to 32 GB. Beyond those sizes, there is more scatter in data extraction due to high file-system load. This is because load on the worker node's file-system can be unpredictable and can affect the data extraction times negatively. Nevertheless, when comparing our extraction tests and processing for the past six months, the predictions by our models (Figure 6.12) correspond to the production runs (Figures 6.15 and 6.16).

Part of the LOFAR SKSP processing is done on shared infrastructure, which requires requesting processing time ahead of time for each grant period. Being able to predict the amount of resources required to process data each grant period is required to make a reliable estimate on what resources to request. These results can be used by other projects sharing the SURFsara GINA cluster to predict their processing time before submitting jobs.

## 6.6 Applications and Conclusions

Our performance model shows that it is possible to predict the processing time and computational resources used by a complex astronomical pipeline. Our results suggest that LOFAR LoTSS processing can be further optimised without sacrificing the

quality of the final product. Additionally, our results are transferable to other scientific pipelines that process LOFAR data with different parameters. Any pipeline that performs gain calibration or application of calibration solutions will benefit from these results.

In order to provide LOFAR processing as a service to scientific users, we need to estimate the processing time for each request. We need this estimation in order to determine whether the user has sufficient resources left in their resource pool. Knowing the performance of the software pipelines as a function of the input parameters will help predict the run time for each request and the resources consumed. Knowing this will make it possible to notify users how long the request will take and how much of their quota will be depleted. It is important to note that while this model is specific to LOFAR processing, the method we detail can be used by other scientific teams in order to predict the computational requirements for their pipelines. Doing this is necessary if large scale scientific processing is to be offered as a service to the wider community.

Finally, a performance model of the LOFAR software will help make predictions on the time and resources needed to process data for other telescopes such as the Square Kilometer Array (SKA). Once operational, the SKA is expected to produce Exabytes of data per year. Processing this data efficiently requires understanding the scalability of the software performance to facilitate scheduling and resource allotment. Overall, we show that our method helps guide algorithm development in radio astronomy, can be used to predict resource usage by complex pipelines and will be promising in optimising data processing by future telescopes.

## 6.A Calibration Solutions for the sky model tests

The output of the calibration step is a data set corrected for direction independent effects, as well as a set of calibration solutions. Figures 6.19 and 6.21 show the calibration solutions for core stations obtained when calibrating with sky models with minimum flux cutoffs of 0.05, 0.3, 0.8 and 1.5Jy. Much like in Figure 6.7, we can see that there is no significant difference between the calibration solutions for these stations. As a note, the naming scheme for LOFAR stations is CS/RS for core/remote stations, three digits for station number, HBA0/HBA1 for High Band antennas and LBA0/LBA1 for low-band antennas. The 0,1 suffixes correspond to sub-arrays in the core stations, which can be correlated separately. Additionally, the CS/RS is replaced with the 2-letter country code for international stations[155].

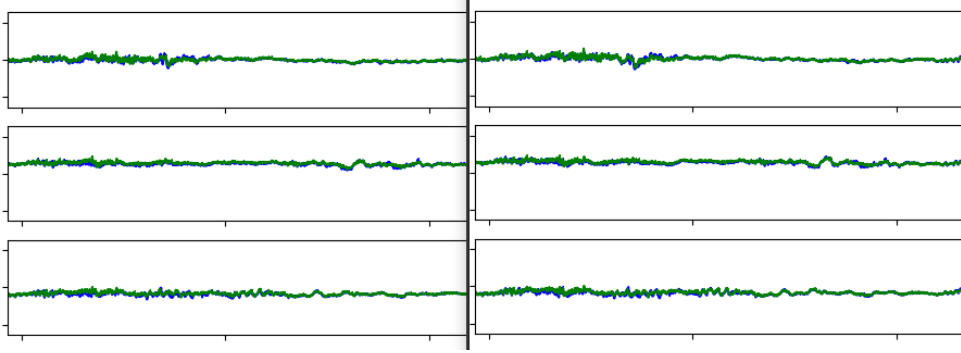


Figure 6.19: The calibration (phase) solutions for the test data set obtained when calibrating with sky models of 0.05 Jy cutoff (left) and 0.3Jy cutoff (right). The data shows the phase solutions for baselines including stations CS003HBA0, CS003HBA1 and CS004HBA0, with respect to the reference station, CS001HBA0. The right solutions were obtained using the production calibration model. We do not see any improvement in results in the left figure, which took twice as long to obtain.

We compare the phase solutions for stations CS032HBA0 and CS003HBA0 and the reference station for two different calibrations in Figure 6.20. We note that this difference is within 0.3 rad for the entire observation. Combined with the results in the other two plots, our results suggest that the calibration solutions do not degrade when calibration is done with a smaller sky model.

## 6.B Parametric model parameters and fit accuracy

In this section, we note the uncertainties to the models fit in Equations 6.1-6.5.

### 6.B.1 Fits quality of run time vs input size model

The models of the processing time vs input size were fit as a linear regression. In this work we present such models for the `gsmcal_solve`, `gsmcal_apply`, `dpp-pconcat`, `predict_ateam` and `ateamcliptar`, the five slowest steps. The resulting models, calculated by the `scipy linregress`[156] routine, are shown in Equation 6.1. We present the  $R^2$  values, P values and standard error below, in Table 6.5.

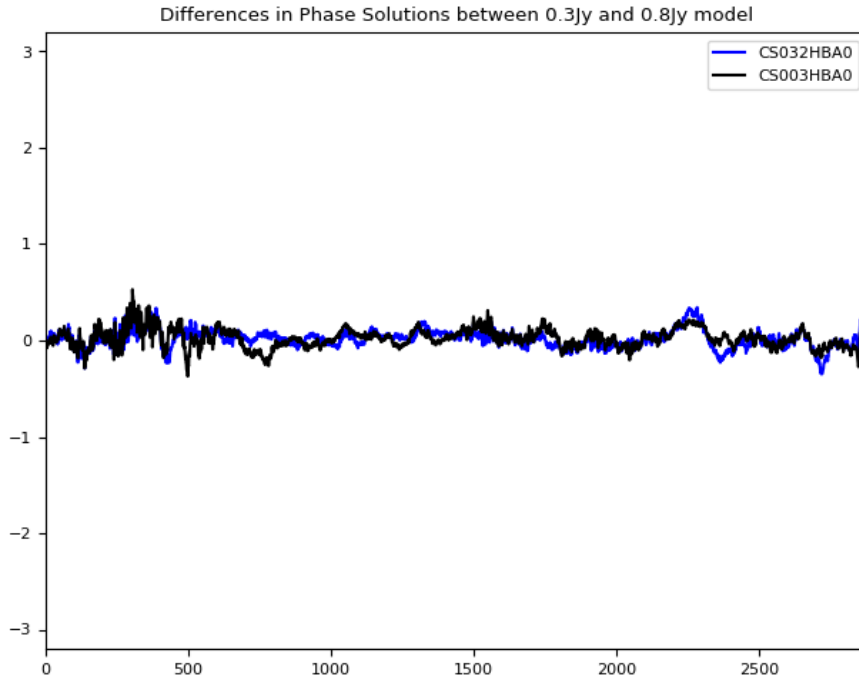


Figure 6.20: Difference of phase solutions between calibrations with the 0.3Jy and 0.8Jy sky models. The solutions for both stations are around zero phase for the duration of the observation.

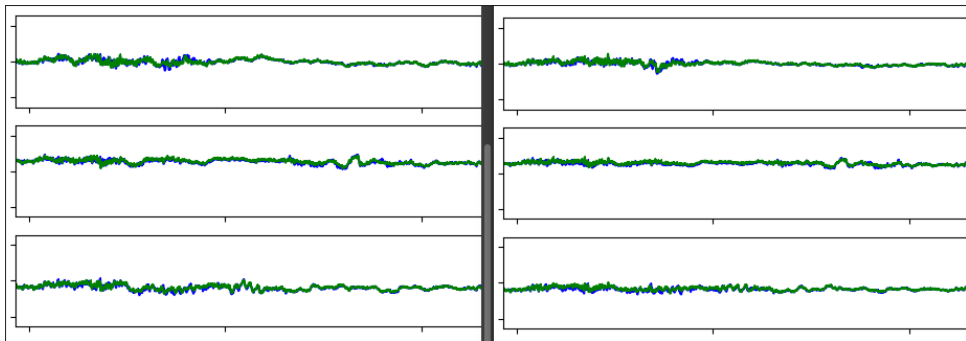


Figure 6.21: The calibration (phase) solutions for the test data set obtained when calibrating with sky models of 0.8 Jy cutoff (left) and 1.5Jy cutoff (right). The data shows the phase solutions for baselines including stations CS003HBA0, CS003HBA1 and CS004HBA0, with respect to the reference station, CS001HBA0. We can see that the calibration solutions shown here are not significantly different than those shown in 6.19, despite taking a fraction of the processing time.

| prefactor<br>step      | $R^2$ | P value                | Standard Er-<br>ror    |
|------------------------|-------|------------------------|------------------------|
| predict_ateam          | 0.996 | 0                      | $1.92 \times 10^{-10}$ |
| ateamclipar            | 0.979 | 0                      | $3.94 \times 10^{-11}$ |
| dpppconcat             | 0.999 | $1.2 \times 10^{-128}$ | $1.78 \times 10^{-10}$ |
| gsmcal_solve<br><=16GB | 0.995 | $3.12 \times 10^{-75}$ | $6.80 \times 10^{-9}$  |
| gsmcal_solve<br>>16GB  | 0.951 | $7.07 \times 10^{-40}$ | $1.58 \times 10^{-8}$  |
| gsmcal_apply           | 0.989 | $5.6 \times 10^{-82}$  | $3.12 \times 10^{-10}$ |

Table 6.5: Fit parameters for the models in Equation 6.1.

$$\begin{bmatrix} 6.83 \times 10^2 & -1.94 \times 10^{-1} \\ -1.94 \times 10^{-1} & 5.81 \times 10^{-5} \end{bmatrix} \quad (\text{B.7})$$

Equation 6.7: The covariance matrix of the parameters in model in Equation 6.2.

### 6.B.2 Fit of run time vs calibration model flux cutoff

The run time vs Flux cutoff model shown in Equation 6.2 is defined by the equation  $y = a \cdot x^{-k}$  and two parameters,  $a$  and  $k$ . The covariance matrix for these two parameters is shown in Equation 6.7. The standard deviation for the fit of the parameters  $a$  and  $k$  is 26.134 and  $7.624 \times 10^{-3}$  respectively.

### 6.B.3 Fit of the NCPU model

The covariance matrix for the fit parameters of equation 6.3,  $a$  and  $k$  in  $y = a + \frac{k}{N}$  are shown in Equation 6.8. The standard deviation of the fits for  $a$  and  $k$  are 13.11 and 48.20 respectively.

$$\begin{bmatrix} 171.94 & -504.11 \\ -504.11 & 2322.95 \end{bmatrix} \quad (\text{B.8})$$

Equation 6.8: The covariance matrix for the parameters for the model predicting run time vs Number of CPUs used, shown in Equation 6.3.

$$\begin{bmatrix} 2.53 \times 10^{-4} & 1.08 \times 10^{-3} \\ 1.08 \times 10^{-3} & 4.60 \times 10^{-3} \end{bmatrix} \quad (\text{B.9})$$

Equation 6.9: The covariance matrix for the parameters for the model for Download and Extract time, shown in Equation 6.5.

#### 6.B.4 Fit for the queuing time model

The statistics of the model fit parameters for the queuing time model (Equation 6.4) are in Table 6.6. The queuing model is fit to the 75<sup>th</sup> percentile of the queuing times for each parameter step. Since this results in a single number for each step, the model's P values are larger than the models from the other sections.

| Value of $\mathcal{N}$ | $R^2$ | P value | Standard Error |
|------------------------|-------|---------|----------------|
| $\mathcal{N} \leq 4$   | 0.382 | 0.381   | 37.086         |
| $\mathcal{N} > 4$      | 0.986 | 0.075   | 86.293         |

Table 6.6: Goodness of fit parameters for the model in Equation 6.4. Since the model is split into two parts, we treat each section as a single linear model.

#### 6.B.5 Fit of the download and extract model

Equation 6.9 shows the covariance matrix for the two parameters  $a$  and  $k$ ,  $y = a \times 10^k$  with the best fit values shown in Equation 6.5. The standard deviations of the fits for  $a$  and  $k$  are 0.016 and 0.068 respectively.

**Original Abstract:**

Data collected by modern radio telescopes are typically many terabytes in size. Testing even simple processing tasks on data of this size is computationally expensive. LOFAR is a modern low-frequency radio telescope that produces petabytes of data per year. The LOFAR radio telescope uses several complex pipelines to process archived data. The development pace of these pipelines is such that their code and parameters can change multiple times per week. To ensure that LOFAR scientific pipelines preserve data quality, we require automated testing and validation of output data. We introduce a method to automate testing of large radio data sets and their accompanying pipelines. Our software is integrated with a leading High-Throughput computing cluster and can scale according to the processing requirements of the scientific pipeline. Using our method, we discover a change in output data quality and track this change to a specific processing software packa, ge.

The contents of this chapter are based on a manuscript to be submitted to Astronomy and Computing.

## 7.1 Introduction

Modern astronomical instruments produce increasingly large sets of data, often in the range of petabytes per year. Scientific insights into this data are obtained typically by individual researchers or small groups, using their custom-tailored processing scripts. These scripts can evolve into a data reduction pipeline which becomes

a standard processing step used to produce scientific products, serving the broader scientific community. Often these scripts evolve into complex scientific pipelines before implementing rigorous checks on the quality of the data products. Moreover, the lack of existing automation infrastructure makes it likely that software updates can change the quality of the processed data without the knowledge of the scientific developers or larger scientific collaborations. In this work, we describe an implementation of an efficient automated system focused on testing complex scientific pipelines for radio astronomy. Our tests cover the initial processing pipeline necessary for producing LOFAR broadband images; however, this implementation can be used by other large scale astronomical telescopes.

The Low Frequency Array (LOFAR) radio telescope is a European low frequency aperture synthesis telescope operated by the Netherlands[77], that observes astronomical objects in the 10MHz-240MHz frequency range. It is designed as a flexible telescope able to serve multiple science cases, such as extensive surveys, transient studies, studies of galactic and extra-galactic sources, radio spectroscopy, as well as long-baseline interferometry.

LOFAR data is stored as a Measurement Set [33] and archived at one of the three Long Term Archive locations<sup>1</sup>. The Measurement Set standard and the high time and frequency resolution of the archived data allow a single observation to serve multiple science cases. Additionally, the standard data structure is supported by several software packages which can manipulate the data in the Measurement Sets according to the scientific goal and include their data products to the Measurement Set.

The scientific data products for each project are produced by a series of transformations of the raw data. These transformations are performed by a set of scripts termed a scientific pipeline. These pipelines are typically written as Python scripts which call various software packages with parameters specific to the pipeline and science case.

One crucial scientific pipeline is the *prefactor*<sup>2</sup> [16] Direction Independent (DI) calibration pipeline. It aims to remove effects from the broadband data that do not depend on direction, such as removing radio interference, flagging bad data and antennas, removing contamination from bright off-axis sources, and calibrating the data against a model of the radio sky [17, 18]. Correcting for these effects is a necessary step in obtaining a scientific image. These steps remove image

---

<sup>1</sup><https://lta.lofar.eu/>

<sup>2</sup><https://github.com/lofar-astron/prefactor/>

artifacts caused by direction-independent effects. The resulting data is processed by a Direction Dependent (DD) Calibration pipeline, which corrects for direction-dependent effects such as ionospheric disturbances.

Scientific pipelines, such as `prefactor`, consist of a set of scripts that use compiled LOFAR software, expected to be installed in the processing environment. Separation of processing scripts and compiled software may result in incompatibility, as they are developed by separate groups. Incompatibilities can often lead to either failure in the pipeline or corrupted data. Often these errors are discovered weeks or months after the introduction of the bug, due to the lack of integrated testing of the pipelines. A further issue is that both the scientific pipelines and the underlying processing software are developed at a high pace. Often, this development pace produces a mismatch between the software and pipelines, creating errors in the scientific images. Automated testing is required to ensure the interoperability of LOFAR software and LOFAR scientific pipelines.

The large data sizes of LOFAR observations pose two challenges to automated testing of scientific software and complex pipelines: (1) the processing time required for continuous testing is significant when compared to the resources available; and (2) no framework exists to automate LOFAR processing. We aim to solve these challenges by integrating automated software quality tests at the Dutch National Grid Infrastructure at SURFsara, Amsterdam.<sup>3</sup> We build on previous advances in automated High Throughput astronomical workflows, such as the integration of a workflow orchestration software (Apache Airflow) with a High Throughput Computing (HTC) cluster[37]. We show that our work overcomes processing challenges, and automates testing of scientific software and pipelines. These advances allow us to detect processing errors early, and to ensure that the scientific pipelines do not degrade the data products.

**Contributions:** The main features of this work are the following:

- Design of a platform to perform integration tests of (multiple) complex scientific pipelines on a High-Throughput architecture.
- A system for automatic versioning of the releases with a possibility of automating deployment and versioning, making data reduction of LOFAR reproducible.
- Methods for performing automated data quality checks. This makes fast development cycles of scientific pipelines possible.

---

<sup>3</sup><https://userinfo.surfsara.nl>

**Outline:** The organisation of this manuscript is as follows: We provide details about the LOFAR use case and corresponding scientific aims in Section 7.2. We describe previous work in Continuous Integration and automation of scientific processing in Section 7.3. We detail the workflow in Section 7.4. We discuss our results and show quality metrics for several regions in Section 7.5. Finally, we conclude and make remarks on future development in Section 7.6.

## 7.2 Background

One of the major Key Science Projects for the LOFAR telescope is the Surveys Key Science Project (SKSP)[157]. The goal of the SKSP project is to create multiple large-scale maps of the northern sky at low frequencies and unmatched angular resolution and sensitivity. To reach this high sensitivity, and to serve multiple science cases, each observation is integrated for roughly eight hours. Observations are stored at a high time and frequency resolution, at 1 second and 2 kHz per sample. This is done to be able to serve multiple scientific goals with the same archived data set. One of the surveys projects is the LOFAR Two-Meter Sky Survey, LoTSS [25]. This survey is expected to produce more than 3000 8 hour observations, each of which is up to 16TB in size. The size and number of the LoTSS data sets pose a significant infrastructure and processing challenges, for which a solution is described in our previous work [35].

Over the past two years, we have processed the bulk of LoTSS data on the GINA high throughput cluster<sup>4</sup> located at SURFsara, the Dutch national High Performance Computing Centre. To help automate this processing, we have built a workflow management system, AGLOW, which is based on Apache Airflow, a workflow orchestration software, able to schedule and execute multiple workflows.

Like Airflow, AGLOW can implement arbitrary workflows encoded in Directed Acyclic Graphs (DAGs). DAGs are a type of directed graph, without internal cycles. A workflow can be encoded in a graph, by having each of its constituent tasks mapping to a node on the graph, and the prerequisites for each task mapping to the incoming edges to that node. Software such as Airflow resolves the requirements for each task, schedules, and executes them efficiently. In production, AGLOW implements the `prefactor` pipeline as a DAG. In AGLOW, we take advantage of the parallelization available on the GINA HTC cluster to process LOFAR data efficiently. Specifically, we implement the `prefactor` scientific pipeline, which consists of two

<sup>4</sup>Specifications of the GINA cluster can be found at <http://docs.surfsaralabs.nl/>.

main steps: calibration of a short calibrator observation and the calibration of the science target. Because of the frequency independence of these steps, calibration of both these observations can be done in parallel.

In software engineering, Continuous Integration is the practice of setting up tests that automatically run as code when a repository is updated. These tests ensure that new code added to the repository do not introduce bugs or errors in the data processed with the code. Because of the complexity of modern scientific pipelines, performing automated tests is necessary to ensure the quality of the output data is preserved. Running these tests on a highly parallel platform, such as the GINA cluster at SURFsara, makes it possible to efficiently test scientific pipelines using large data sets and ensure the quality of the data products is preserved despite the fast development pace of software.

The LOFAR Surveys project distributes its software in the form of pre-built images using Singularity[158]. Singularity is a containerization software allowing the user to build and distribute software images that can be used by an unprivileged user to access the contained software. Singularity is successfully used at multiple university clusters, and is already installed and tested at the GINA cluster. Before safely deploying software images to multiple scientific teams, we need to verify that the compiled software is compatible with scientific pipelines, and do not degrade the quality of the final scientific image.

## 7.3 Related Work

Continuous Integration (CI) [159] and Continuous Delivery (CD)[160] are concepts that have been used by software developers over the past decade and a half. Continuous Integration allows for large teams to collaborate on software without worrying about introducing bugs in their development process. This method relies on unit tests and integration tests being shipped with the software, and includes a system that runs such tests every time code is updated. Continuous Delivery builds on this process to automatically validate output products in a production environment, and release a working software package.

Large scientific teams have begun introducing CI workflows in their development cycle with many previous successes [161–164] in various fields from genetics to telescope control. These works show the power of CI systems to enable rapid software iteration in an academic environment.

Commercial integration services such as TravisCI[165], CircleCI[166], and GitLabCI[167] are widely used in industry and combined with GitHub or GitLab for source control. Additionally, automation suites such as Jenkins[168] are typically used to automate CI pipelines on dedicated infrastructure. While Jenkins is a fully-featured automation suite, it needs to be installed and managed by a user with elevated privileges. These advances can also be seen in the new drive to provide a Platform As A Service cloud infrastructure to scientists [169]. Our own work creating a distributed platform for LOFAR processing [50] describes solutions to the technical requirements of large scale LOFAR processing.

Apache Airflow<sup>5</sup> is an Apache workflow orchestration software built in Python. It allows building and scheduling generic workflows, and in previous work we implemented the full LOFAR prefactor pipeline running on a distributed infrastructure using Airflow (which we named AGLOW) [37]. We deploy this software in user-space, allowing us to easily build and manage LOFAR workflows without requiring elevated privileges. AGLOW has been proven useful in running LOFAR processing on a shared cluster.

With the advancements in containerization, more and more groups have begun distributing scientific software in pre-built containers. Docker is a common way to package, version, and distribute scientific software, and its use has ensured that research is easily reproducible [170]. A competing, open standard is Singularity [158], built by Berkley National Lab. Unlike Docker, Singularity images can be used without elevated privileges and can be distributed as an executable file. Due to its flexibility, many scientific groups now use Singularity in their scientific workflows [126, 171–173]. Because we use shared infrastructure to process LOFAR data, we do not have exclusive, administrative access to our processing machines. As such, we opt to use Singularity for software distribution. Because Singularity images are built from text-file recipes, we store these recipes in a GitHub repository which makes it easy to determine the changes between software images.

## 7.4 Automated testing with AGLOW

We use our previous success in automating complex LOFAR pipelines with AGLOW and build a Continuous Integration workflow. We deploy this workflow on the shared infrastructure available at SURFsara. Thus, our tests run on

---

<sup>5</sup><http://airflow.apache.org/>

the same cluster used for the LOFAR SKSP processing, which ensures compatibility of deployed software with our processing infrastructure. The source code of our CI workflows is located in the AGLOW software repository located at <https://github.com/apmechev/AGLOW/tree/master/AGLOW/airflow/dags>.

Our software tracks the version history of one or several scientific pipelines using the API provided by GitHub. Using this API, we check for new updates to the software pipeline(s) and trigger the test workflow associated with each pipeline. In a similar manner, the build scripts for our Singularity images are stored on GitHub. In the case that these images have been updated, we trigger the test workflow for all scientific pipelines ensuring that we verify that a change in the software does not corrupt the data produced by any of the scientific pipelines.

### 7.4.1 Distributed CI Workflow

Figure 7.1 shows a run of the CI workflow in progress. The workflow checks the latest commit of both the `prefactor` repository and software image and compares it with the date of the last completed run of both. If testing is required, the workflow sends a request to bring the test data (observations of a calibrator source) to disk. Once the calibrator data is available, the calibrator part of the `prefactor` pipeline is tested.

Once the calibration is completed, the same steps are repeated with the `prefactor` scripts for the science target. Upon successful completion, the ‘`move_results`’ task uploads the software image to our distribution repository. We currently distribute software on a WebDAV[174] server hosted by SURFsara, making LOFAR software easily available to users worldwide.

### 7.4.2 Nightly Builds

Because running the `prefactor` pipeline on our test data set (described in the Appendix) takes several hours and several hundred CPU-hours, we only test the software on a nightly basis. We run this nightly build process at midnight every day. If the tests succeed, we upload the new software images, scientific images, and an archive of the pipeline scripts to a deployment area identified by the date of the run. A diagram of the test process is shown in Figure 7.3. In this work, we compare scientific data produced with our CI pipeline over the span of five months. The results presented in Section 7.5 indicate the utility of automated tests of LOFAR data processing.

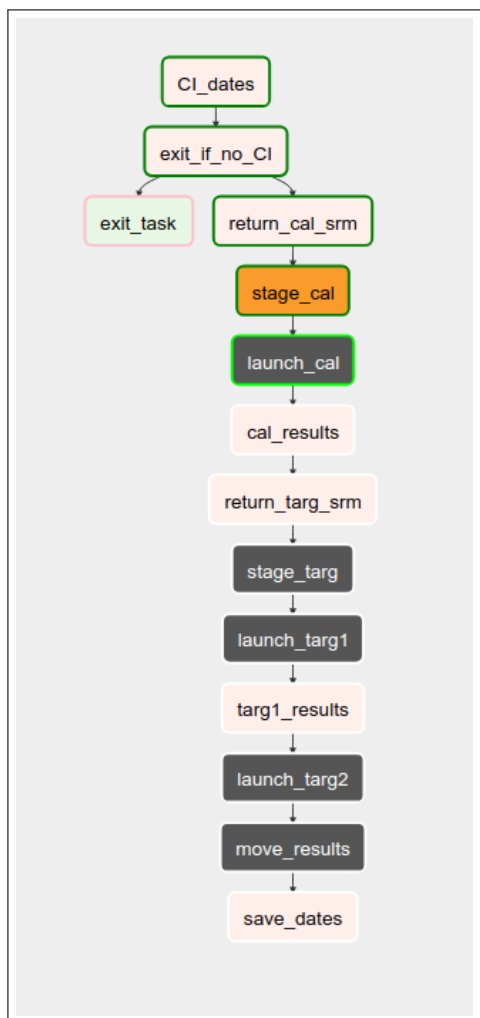


Figure 7.1: The workflow responsible for testing the `prefactor` pipeline as well as the current LOFAR software image. The tasks in this workflow check if the software image or the scientific pipelines have been updated. The workflow stops if there have been no updates. In the case that tests need to be run, the data is requested from the LOFAR archive and the `prefactor` steps are executed on the test data. The three processing steps are the calibrator calibration ('`launch_cal`'), Target pre-processing ('`launch_targ1`') and Target calibration ('`launch_targ2`'). In between each step, there are checks whether the required data products have been produced.

### 7.4.3 Testing of Software Images

We build the LOFAR software on Singularity-Hub by integrating a GitHub repository containing the build scripts<sup>6</sup> with the Singularity-Hub builders, and storing the completed image directly in Singularity-Hub<sup>7</sup>. Singularity-Hub allows multiple versions of an image to be hosted and differentiates releases by their MD5 hashes. We use the ‘frozen’ version of a container as a stable release and the most recent version as a test release.

In addition to the `prefactor` pipeline, we also run the unit tests of the `numpy`[175], `scipy`[156] and `astropy`[176] scientific libraries. We import and run the tests module for each of these libraries. Testing the underlying libraries is necessary to ensure that we maintain the numerical accuracy of all LOFAR processing.

### 7.4.4 Testing of Scientific Pipelines

To test scientific scripts, we use the latest commit of the GitHub repository. These scripts take as input a set of parameters that can be modified by the user. We use the same parameters as in the LoTSS processing, to ensure that the image quality obtained during our tests will match the quality of our production run. Once the data is processed, our tools store the intermediate products with a time-stamp. This allows future comparison between data processed with different versions of the software and scripts.

## 7.5 Results

Here, we present our results, implementing the first automated testing of LOFAR data processing on a High Throughput infrastructure. We discuss our solutions to the challenges of automatically processing large volumes of data, as well as discoveries made from the analysis of the resulting data. The infrastructure described is suited for projects using complex, computationally intensive scientific pipelines interested in preserving the quality of processed data even at a high development pace. The tools we have built are general in that they can test any complex processing pipeline as long as it is hosted on a git server. Specifically, it can accelerate massively parallel pipelines and handle data sets of up to several terabytes.

---

<sup>6</sup>located at <https://github.com/tikk3r/lofar-grid-hpccloud>

<sup>7</sup>Images are hosted at <http://singularity-hub.org/collections/1999>

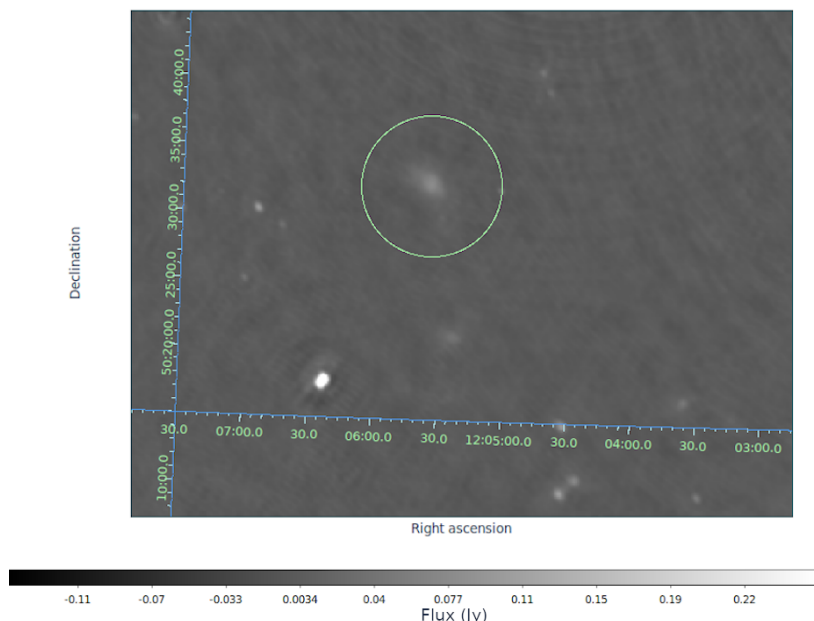


Figure 7.2: An image of the test set created on 2019-04-23. We sample the diffuse region selected to get statistics in Figure 7.5. The image contains some calibration artifacts that will be removed by the Direction Dependent Calibration that follows the `prefactor` pipeline. The images produced between March 2019 and August 2019 do not appear qualitatively different, however the measurements shown in Figures 7.5 show there are some changes in the resulting data.

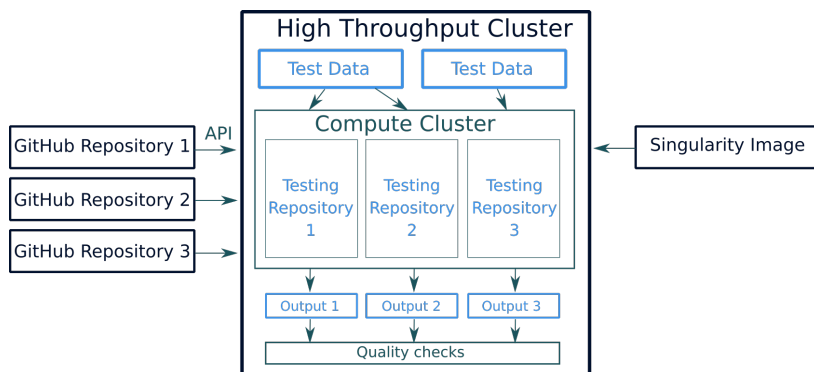


Figure 7.3: Figure showing three scientific pipelines, hosted on GitHub tested on a High Throughput infrastructure. We access GitHub through the GitHub REST API, store test data and results on dCache [130] storage at the site of our processing infrastructure. Scientists can use their custom scripts to test the quality of the final data products.

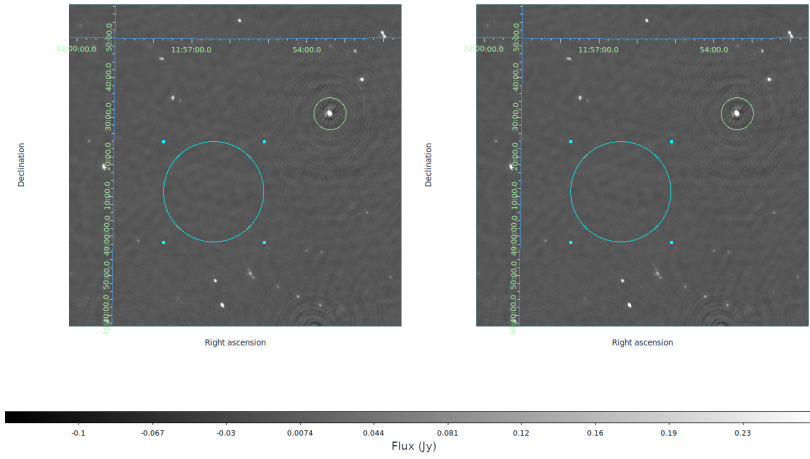


Figure 7.4: Images created by the CI runs from 2019-03-29 (left column) and 2019-04-23 (right column). The green region shows a bright source of  $\sim 6$  Jansky, which we use to validate the extracted flux for point sources. The cyan region is an empty part of the sky that we use to obtain noise measurements of the image. We show measurements such as the integrated, peak flux as well as the RMS noise of the region.

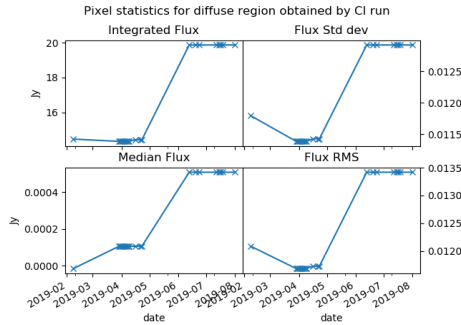


Figure 7.5: Pixel statistics for the region circled in Figure 7.2 obtained from calibrated images produced by the AGLOW CI runs from 2019-03-29 to 2019-08-01. The sharp decrease in data quality corresponds to an update of the underlying software image. We confirm this by running 2019-04-16 pipeline with the image from 2019-08-20 and obtain the same quality as with the 08-01 pipeline. The plots of the other scientifically important regions, the bright source and source-free background, show a similar change. Scientists can use these plots to quickly detect changes in image quality.

### 7.5.1 Pipeline Automation and Implementation

The goal of this work is to present the first automated solution for automated testing of sophisticated scientific software and pipelines. We implement the complex LOFAR `prefactor` pipeline as a first use case. Performing a full test of this pipeline is both computationally expensive and challenging to automate, which is why we leverage previous successes in large scale distributed automation of LOFAR processing to enable this study.

### 7.5.2 Data Quality

Astronomers use multiple methods when comparing data quality of final images. While the results from the `prefactor` pipeline need further processing to obtain a final scientific image, the intermediate results produced by our tests can still be compared across different software versions. This comparison will give early bounds on the final image quality and is an essential part of producing a high fidelity final image.

The primary way to compare data processed by `prefactor` is by eye, however quantitative metrics such as the background noise and integrated flux around sources are also used to compare results. In addition to images, the resulting data includes calibration solutions. The statistics, such as the mean phase, phase variance, and quartile range of these solutions are useful scientists to determine calibration accuracy. This work presents a method to automatically track these metrics and report any deviations.

One large deviation in the metrics was introduced by updates to software or pipelines, appears as a jump in Figure 7.5. These figures suggest that an update in software between April and June 2019 has affected the scientific results. We confirm this result by processing the April 2019 version of `prefactor` with the August software image. The results of this test are the same as the automated runs from August, which rules out changes in `prefactor` as the cause. When comparing the GitHub repository containing the recipes of the software images, we discover that the software performing ionospheric fitting has been updated between April and August. Without our automated testing, this change in image quality will likely not have been discovered, and the underlying software change would have been difficult to trace.

We show the images obtained from our automated runs in Figures 7.2 and 7.4. The former figure shows 13 CI runs, highlighting a region of diffuse emission,

while the latter shows a bright point source and an empty region for calculating noise statistics. Using statistics from all three regions, we can track the ability for the pipeline to image diffuse sources, compact sources, and retrieve faint sources. All three types of sources are important for different radio astronomy science cases. Measurements around the diffuse source indicate how accurately the software extracts fainter, diffuse sources. We have also chosen a bright point source (the quasar 4C 49.22) in our field to show the ability for `prefactor` to retrieve bright point source fluxes. Likewise, the change in software between April and August 2019 has increased the flux extracted around this point source, as well as the noise in around it.

## 7.6 Discussion and Conclusions

In this work, we present a method to automatically test complex scientific pipelines on large datasets leveraging a High Throughput infrastructure.

During the first five months of `prefactor` Continuous Integration tests, we produced more than 30 data sets, each processed with a different version of the `prefactor` pipeline, using two different software images. We automatically created images of these data sets and compared different metrics for each image. From those results, we could detect the effects of changing processing parameters such as the time and frequency resolution of the data. We discovered that an update in the LOFAR software, specifically a package named `LoSoTo`, in May 2019 has led to a significant change in image quality. `LoSoTo`<sup>8</sup> is a tool for processing LOFAR solutions and for fitting and removing ionospheric effects.

Continuous automated testing is crucial for determining whether changes in software or processing scripts and parameters will affect the quality of the scientific data. In our case, the quality metrics we track are the noise levels and peak/integrated flux levels. Our framework supports arbitrary processing scripts, and thus allows scientists to define and deploy their metrics of interests. Running these tests every night makes it possible to create the first time series of data quality for LOFAR data. These results can be used to notify astronomers when the quality of the scientific products changes or degrades significantly.

Additionally, the `AGLOW` software allows for multiple scientific pipelines to be tested concurrently. Each pipeline will produce its own set of images and qual-

<sup>8</sup><https://github.com/revoltek/losoto>

ity metrics, which are used by astronomers to verify their scripts. As all LOFAR scientific pipelines are hosted on GitHub, they are easy to integrate with our software. In the event of a change in data quality, developers can easily track down the cause by checking the commit history of the scientific pipeline or software image recipe.

### 7.6.1 Conclusions

In this chapter, we describe the first automated, high throughput testing of LOFAR scientific pipelines. We can test LOFAR pipelines and software, create radio images, upload and version software, and report image quality and statistics automatically. Our work shows that it is easy to detect changes in image quality caused by software, algorithm, or parameter changes in the `prefactor` scientific pipeline. Using our results, astronomers can gain insights and help detect degradation of image quality caused by software updates. We support this claim by detecting a noticeable decrease in image quality caused by a software update between April and June 2019. Without our automated tests, this degradation will likely not have been detected or reported.

The flexibility of our solution allows testing commits of scientific pipelines in the past, comparing the data products from historical versions of the software to the current version. Moreover, we can see that most changes in the pipeline do not have a noticeable effect on the data quality, meaning that the scientific quality of the final images is expected to be stable across most commits. Having a time series of image statistics makes it easier to detect, understand, and fix code that leads to significant changes in image quality.

Our successes testing the `prefactor` pipeline was the first case of Continuous Integration used to verify scientific data quality for the LOFAR surveys. The extensibility of the AGLOW system makes it easy to add further pipelines or quality checks to ensure the high development pace does not result in degradation of the scientific products. The work described in this chapter is designed to help large astronomical collaborations ensure high data quality of scientific products, even when confronting considerable data sizes.

## APPENDIX

We test our processing pipelines on a standard LOFAR surveys observation designated L229587 in the direction of the HetDex field[177] in the direction

(11h55m41.282, +049d44m52.908). The data is stored at the SURFsara Long-Term Archive location and can be accessed through the following URI:

```
srm://srm.grid.sara.nl:8443/pnfs/grid.sara.nl/data/
lofar/ops/projects/lc2_038/229587/
```

To minimize processing time and resources, we only use a fraction of the entire frequency bandwidth corresponding to Subbands 100-110. This corresponds to a frequency between 139.844 and 141.602 MHz. The data was observed on 28-May-2014 between 15:00 and 23:00. The initial data was 150GB and the final averaged data set was 2.3 GB.

We make images with the following software and parameters:

```
wsclean -size 2560 1080 -maxuv-l 7000 -baseline -
averaging 5.34930608721 -local-rms-method rms-with-
min -mgain 0.8 -auto-mask 3.3 -pol I -weighting-
rank-filter 3 -auto-threshold 0.5 -j 5 -local-rms-
window 50 -mem 20 -weight briggs 0.0 -scale 0.00208
-niter 5000 -no-update-model-required
```



As astronomical observatories produce ever-growing data-sets, the processing challenges for these data will continue to increase. Extensive astronomical surveys, expected to create petabytes of data, can no longer be processed on a single machine or small dedicated clusters at scientific institutions. Serving the scientific requirements of these surveys will require large scale distributed processing.

CERN's World-Wide computing Grid provides sufficient resources for such projects; however, its focus is on distributed Monte-Carlo simulations. This high-throughput infrastructure offers opportunities for parallel processing of radio astronomy data sets. To take advantage of these resources and implement complex astronomical workflows to a grid-like environment requires a framework to distribute and monitor jobs. Furthermore, processing and re-processing thousands of observations efficiently requires workflow orchestration software. We aim to enable the 30+ petabyte LOFAR Two-meter Sky Survey (LoTSS) by combining high throughput processing infrastructure with modern workflow orchestration software.

## 8.1 Summary of Thesis Contributions

The work in this thesis focuses on software built to accelerate, parallelize, and automate LOFAR processing as well as the insights obtained into large scale processing of LOFAR data. We have built a scalability model which we use to understand the performance of LOFAR broadband processing pipelines. Our model brings novel insights into the limits of our current pipelines, as well as suggestions to improve processing throughput.

We have built a platform for processing a radio astronomy data on a heterogeneous, distributed infrastructure. We exploit the data-level parallelism of com-

putationally intensive processing tasks, and our work makes several LOFAR scientific projects possible. Additionally, our insights into distributed execution of complex pipelines are crucial for enabling sizeable astronomical surveys. We expect distributed processing will become an increasingly important paradigm in astronomy.

Finally, we created an automated workflow system with the goal to automatically produce high fidelity images from LOFAR observations. This system was a successful integration of industry software into radio astronomy, one of the goals of this thesis and its associated grant. Bringing open-source tools used in industry is crucial to keeping long-lived scientific projects maintainable and productive. Our results were an important step in enabling high throughput, automated processing of LOFAR scientific workflows.

Our advances in understanding LOFAR processing inefficiencies, exploiting data-level parallelism, and automating workflows are important steps to modernizing LOFAR scientific processing. The lessons learned in this work can be directly applied in other scientific fields that need to process data at overwhelming rates.

## 8.2 Answers to Research Questions

***Research Question 1:** How can we use a distributed shared infrastructure for efficient LOFAR data processing?*

In Chapters 2 and 3, we present our results enabling massively distributed processing of LOFAR data. We describe the underlying platform, inherited from the High Energy Physics community and the modifications to these tools that were required to host sophisticated processing software. We describe these modifications and discuss the resulting increase in throughput. Finally, we estimate the time saved by parallelizing LOFAR data processing. The work described in these chapters is essential to producing scientific data sets at a high rate, particularly considering the high data rates produced by LOFAR.

***Research Question 2:** How can we build software to effortlessly accelerate complex pipelines for radio astronomy?*

Chapters 5 and 7 present our work on parallelizing LOFAR scientific pipelines on a distributed shared infrastructures. We integrate a mature workflow orchestration package with distributed LOFAR processing. We discuss the need for this orchestration, as well as the abilities to support additional complex pipelines.

As an example application, we build a Continuous Integration pipeline tasked with verifying and validating the initial steps of LOFAR processing.

**Research Question 3:** *Can we automatically collect performance information during massively distributed processing and predict run times for future data sets?*

Chapters 4 and 6 describe a performance monitoring suite for LOFAR data and our scalability model for LOFAR processing. When running massively distributed processing, scientists are unable to monitor the performance of the underlying software. Collecting these statistics is necessary for understanding processing inefficiencies and suggest ways to accelerate data processing. Performance data can also be used to understand the effect of processing parameters on the resource usage of complex pipelines. We study this in detail, building a model that can be used to understand the scalability of multiple processing steps. This model shows the limitations imposed by available processing resources as well as suggestions on decreasing processing time without sacrificing scientific data quality.

## 8.3 Limitations

Using the software described, the LOFAR Surveys team was able to process several petabytes of archived data and produce scientific quality images. Despite the successes of the project, several issues occasionally impede data processing and prevent rapid deployment of software pipelines.

High throughput processing of LOFAR data requires the initial processing steps to be performed at the data archive locations. While deploying new versions of the LOFAR software and pipelines at SUFRsara is straightforward, the same is not true for other LTA locations. LOFAR data needs to be processed at LTA locations that do not support any modern containerization software nor other software distribution methods. This makes deploying new software is difficult and time-consuming. Additionally, orchestrating jobs at these sites requires additional integration with our job monitoring tools due to lack of internet access from some HTC clusters. Integrating locations not suited for large scale distributed processing is an ongoing challenge for the LoTSS survey and other LOFAR projects.

A further limitation is the authentication of our processing software and the authentication requirements for data access. Distributed processing of large data sets requires having access to the data at multiple locations, however the intermediate products produced by our workflows are not public and require an active x.509

certificate authorised by the LOFAR SKSP project. The current workaround to this limitation is to maintain active certificates at each processing location. Upcoming features of the dCache storage system, bearer tokens called macaroons, will make it possible to overcome this limitation.

Finally, our current software distribution does not assign long-term version numbers to software images and scripts, nor is there a way to store these images or cite them in related papers. Implementing proper versioning is crucial to not only making data processing easily reproducible, but also make it possible to recognize the effort put into building and distributing software images. Overcoming these limitations will enable FAIR science with LOFAR data[178].

## 8.4 Future Work

This work focuses on the substantial gains possible by parallelization of LOFAR data processing. We take in mind the complexity of our processing workflows, the full range of scientific pipelines and the heterogeneous nature of the underlying infrastructure. Because of these factors, a wide range of astronomical pipelines can use the software presented in this work. Future automation of the LoTSS processing requires deciding on data quality requirements at each step and automated re-processing strategies in case a data quality check fails. Implementing intelligent re-processing strategies will reduce the human supervision currently necessary to provide high-quality large-scale surveys such as LoTSS.

Scientific projects with significant data rates such as Gaia and LSST provide users with an integrated environment to efficiently process archived observations. Having such an environment is a necessary step to gain fast and easy to gain insights into LOFAR data. This work presents a method enabling scientists to incorporate processing hosted at scientific institutions and cloud providers to scale scientific processing horizontally.

One application for such large scale distributed processing is the Square kilometre Array. The Square Kilometer Array, (SKA) is a planned aperture synthesis radio telescope expected to have a total collecting area of one square kilometre. It is expected to produce more than 160 TB per day [179], data which needs to be processed. Scaling our tools to SKA-size processing requires a federation of clusters able to handle a high throughput workload. Nevertheless, as the SKA data processing will use different software tools than LOFAR, further study is needed on the optimal processing strategy for each of the many SKA science projects.

In recent years, academia has begun focusing on ease of access and reproducibility of science. Science done with cutting edge instruments, such as LOFAR, tends to be time-consuming to reproduce. Barriers such as setting up a working software environment and downloading massive data-sets prevent scientists from quickly and easily reproducing the results of their peers. These barriers make it difficult to verify the accuracy of discoveries and need to be overcome in order to make astronomy more honest and transparent. Automatically building, testing, versioning and releasing docker and singularity images is critical to making science in radio astronomy easily reproducible. Creating a LOFAR science portal and integrating our tools and software images with this portal is crucial to making radio astronomy both more accessible and more authoritative.



## Acknowledgements

Hard to believe that it was four whole years ago that I returned to Leiden to pursue a Ph.D. at the Sterrewacht. Over these years, I was lucky to make so many colourful friends and vivid memories, all of which made the 'Winter' weather a little more bearable. I'm not sure a few short pages can summarize the beautiful people I've met over the past four years, but I'll have to try nevertheless.

First, I'd like to thank the diverse and energetic LOFAR group. Josh, who always has a new intriguing topic to discuss, and is overflowing with ambitious ideas. Jit, who is that friend that every friend needs. It's no wonder he's the go-to choice for a paronymph: he's reliable and easy-going, and that's no secret! Kim and Pedro, probably the first real users of my software, must be commended on their patience and goodwill. Frits, with whom we've shared too many moments complaining about software over beer (or coffee, or tea, or over messenger), I wish you best of luck with the heavy mantle of making sure our software works. Gabriella, Andrea, and Francesco grouped together for no particular reason at all, running into you randomly (or pseudo-randomly) in Bologna was a pleasant surprise. You'd be happy to know that you'll always be associated in my mind with quality food and quality time together. Federica, Amanda, and Leah, the renowned LOFAR ladies, it is always a pleasure running into you at our conferences and workshops, both in the Netherlands and abroad. I sincerely wish you the utmost success in your careers. Carole, I had lots of fun on the many occasions we spent together. Thanks for being enthusiastic enough to join me at a beer festival just before a long week of LOFAR workshoping. Wendy, Martijn, Reinout, and Rafael, it's a pleasure bumping into you and riffing off some ideas in the KL, or just around the Sterrewacht. Aayush, we've had some good chats together, I hope Rome is treating you well! Nicolina, my half-Bulgarian sista, I miss our workshops and conferences, hope we run into each other again.

To my slav friends, it was a pleasure squatting with you. Although we only had a couple of slav dinner nights together, they were a complete success! Matus, you have been the best neighbour and most reliable friend one could wish for. I wish you all the success and hope we still have beers at Plantage every once in a while. Give the bunnies a hug from me! Lýdia, thank you so much for your kind words and all the chats we had together. Andrej, you are already sorely missed in our Friday borrels. Best of luck in Germany and keep on dancing! Liuba, it was pleasure sailing with a seasoned sailor as you. The stories from your sailing trip still inspire me to one day drop everything and sail the seven seas (yarr). Aleksandrina, it's always great seeing you at LEO events! It was you who introduced me to LEO, and it looked so fun I decided to be part of the board myself! Vanja, the dishes you bring to our slavic dinners were all exceptional! It's been lots of fun to have someone that you can practically understand in your native tongue. Michał, Łukasz, and Aleksandra, it's always a party when you're around! Thanks for bringing the Polish atmosphere with you, and of course thanks for the pierogi.

To all borrel peeps, thanks for all our chats on Fridays! Merrel, Arthur, and Jeroen, thanks for bringing the necessary Dutchness to our borrels. Dilovan, hope you keep channeling your inner Mallory Archer and that the future borrel committees treat you with quality wine! Mantas, keep dropping sick beats at the sterrewacht parties! Anna, u

w0t m8? I'm flattered that you take me as a fashion authority, I'll pretend that I actually am. Sanjana, hope one day you are able to adopt a little duckling. Let me know if you hear about those prestigi/ous post-docs at McDonalds, they're hard to come across I hear. Fraser, Alex, , Maaïke and Patrick, I'll miss having the occasional Canadian corner at our borrel. Marta, it was fun chatting with you over the garlic puzzle. Kirsty, keep being a cool cat (a cool koala? a kooala? Let's go with that). Robin and Malvika, hope you enjoy your time in Leiden as much as I have, and best of luck in future endeavours!

To the big Dipper, thanks for organizing our weekly celebration of a week well spent. I will always remember that there's a slight possibility that bees 🐝 are just mad 🤡 bee🐝cause someone cut✂️them off in traffic🚗🚗🚗🚗🚗🚗. I hope I don't turn into a bee in mid acknow-

Our agraphia group was a fun little exercise and really helped me regiment my writing. Thanks, Francisca, for setting it up. It was fun doing our bi-weekly calls with you Łukasz and Ylva. We made so much progress together! To all my Astronomy On Tap compatriots, it was plenty fun running the events together. Liz and Vincent, thanks for putting in all the work to organize these events, your work really made the events memorable. Fran, Carmen, and Themiya, your hard work makes the events run smoothly. Maaïke, my partner in crime, filming is always more fun when you're around. Best of luck heading the video team when I'm gone! Alex, we gotta catch another Rye and ginger at de Burcht. No Canada Dry, but we'll make do. To the rest of the LOFAR e-infra group, Natalie, Coen, Raymond, Zheng and Tim, Thanks for all the help, ideas and support, both technical and personal. Mark, Maria, Kelly, and Laura, I will miss running into you at API, the occasional party or LOFAR conference.

To the rest of the Observatory, it was my pleasure spending the last four years with such a fun company. Tommaso, Maria Christina, Gabriella, Dario, and Elenora, I'll miss the Italian atmosphere you bring. Emanuelle, I'll also miss the atmosphere you bring (hard to believe, eh?) Marina, Christos and Turgay, it was great chatting with you over coffee and tea.

To the Soccorro pals, Hans, Laura, Fateme, Maja, and Jackie, I will miss our shenanigans that one summer in New Mexico. My LEO fam, Nelli, Redmar, Antonis, Noemie, and Garnet, I had such a fun time organizing events, having our meetings and nights out, and just hanging out with you all! There's way too many memories to mention but I'll always treasure the feeling of excitement that came with the first and feeling of accomplishment that came with our last party. To Lýdia, Laura, Sanne and Julliette, our new LEO ladies, keep up the hard work!

The Science Club Crew, I'll miss our Thursday beers and discussions. Sayan, MJ, and Sander, thanks for making Thursday afternoons more enjoyable. Santiago and Pablo, I miss running into you at the S-club! The Lemmy's crew, feat. Nick and Dionne, I always have a blast having beers with you, whether at Lemmy's, Nick and Dionne's place or just around town. Pranav, you always make for a delightful conversation, let's hope we can play some tennis together soon! Hreinn, Sayan, and Sasha, I feel so much more enriched after our discussions. Francesca, it's a delight when you make it to our evenings. I promise, we're more fun than the lab! Rob, we should catch another football match together over pints sometime! Alex and Eliza, I had lots of fun at the house parties you threw! We need another flipcup night :) Jelle, thanks for the help reading

---

over my summary, for always being there in our group and for being so proactive with meetups, organizing nights out and being a crucial part to our social group. Christina and Lorenzo, it's a blast chatting with you about neuroscience or skiing respectively. Meropi, we miss you dearly in Leiden, hope you drop by soon.

To the rest of the many friends I've made in Leiden, Elena, Vincenzo, Helena, Chris and Al, hope your memories of this place are as fun and plentiful as mine. There are too numerous occasions to recount, however all my memories have been good ones, and I hope the same goes for you all. To Cassidy, Angela, Kyle, Sam, Aarya, and Ellert, I had such a blast at Python in Astronomy! Hope we can have a reunion soon. Ana, Jels, Skralan, Christime, Dylan, Joaquin, Viviana, Mónica, Jorge, Boruta, Zanečka, I miss you all and the great times we had in Leuven. We need to resume our kot-crawls even though the kots are slowly turning into houses and the crawling turning into a full day's travel.

To my supervisors across both sides of the street, Huub and Huib in Astronomy and Aske at LIACS, I am sincerely grateful for the kind guidance and opportunity to explore two very different, but deeply engaging fields. To my supportive family, on both sides of the Atlantic, thank you so much for being there for me and believing in me. Your support, encouragement and kind words are what made this work possible. My paranymphs, Ivan and Rossella, I'm proud to have two such loyal, kind-hearted, and inspiring paranymphs by my side during my defense. Thanks for being here for me.

Finally, to my love, Rossella. My feelings for you can fill an even larger book than this, and that's even without any plots, tables, and figures. But I'll keep it to this: being with you makes life indefinitely easier. I cherish every day we have together, whether sunny or (as is the weather around here) rainy. You are the one person for whom I strive to be the best man I can be, and for that, I cannot thank you enough.



## Samenvatting

---

Wanneer we naar de nachtelijke hemel kijken, observeren we het universum door zichtbaar licht. We kunnen sterren, planeten, nevels en misschien zelfs andere sterrenstelsels zien. De vroegste astronomen bestudeerden het Universum met hun ogen, uiteindelijk geholpen door spiegels en lenzen in de vorm van telescopen en verrekijkers. In de 18e en 19e eeuw hebben natuurkundigen het volledige elektromagnetische spectrum ontrafeld. Naast zichtbaar licht, konden astronomen nu het universum bestuderen met behulp van de infrarood-, ultraviolette, microgolf- en radiodelen van het spectrum. Licht uit elk deel van het spectrum draagt verschillende informatie met zich mee en kan worden gebruikt om specifieke fysische fenomenen te bestuderen. Het nadeel van deze breedte van informatie is dat elk deel van het spectrum zijn eigen specifieke detectoren nodig heeft, en vaak hele telescopen. Toch onthult elk deel van het spectrum een deel van het universum dat voor ons verborgen is in alle andere golflengten.

Radioastronomie werd geboren in de jaren 1930 met de experimenten van Karl Jansky met een gerichte 30-meter radio-antenne. Hiermee kon hij onweersbuien detecteren, de magnetosfeer van de Zon, maar ook een vreemde onbekende bron in het centrum van onze sterrenstelsels. Deze experimenten bewezen dat laagfrequente radiogolven konden worden gebruikt om het verre heelal te bestuderen. In de jaren 1940 werden radarontvangers ontwikkeld als een kwestie van nationale veiligheid tijdens de Tweede Wereldoorlog. Ze worden beschouwd als de belangrijkste reden voor de overwinning op de Luftwaffe in de Slag om Engeland. Na het einde van de oorlog werd een deel van de radarhardware op de ruimte gericht.

Terwijl de nieuwe antennes veel gevoeliger waren voor radiofrequenties, hebben radiotelescopen een fundamentele beperking in termen van resolutie. Omdat licht zich als een golf gedraagt, vooral bij lange golflengten, wordt het beperkt door de diffractielimiet. De diffractielimiet verbindt het kleinste object dat

ruimtelijk kan worden onderscheiden door een telescoop met de golflengte van het licht en de diameter van de telescoop. Als we een radiotelescoop van 100 m nemen, heeft deze de onderscheidingsvermogen die gelijkwaardig is aan die van een zichtbare telescoop met een spiegel met een diameter van 0,5 mm. Dit zou niet al te schokkend moeten zijn: radiogolven zijn meer dan 10 miljoen keer langer dan het licht dat we waarnemen, dus onze telescopen moeten natuurlijk 10 miljoen keer groter zijn.

Om de resolutie van de Hubble-ruimtetelescoop te evenaren, zou de spiegel van een radiotelescoop op 150 MHz een diameter van  $\sim 1000$  km nodig hebben. Terwijl ingenieurs nog steeds bezig zijn met het maken van radarchotels van dit formaat, hebben astronomen zich tot computers gewend om de hoekresolutie van radiotelescopen te vergroten. Met behulp van de golfeigenschappen van licht kunnen radioastronomen de timing synchroniseren van de gegevens die zijn waargenomen op meerdere antennes, gescheiden door tientallen, honderden of duizenden kilometers. Met aanzienlijk rekenwerk kunnen gegevens van deze ver uiteen liggende antennes worden gecombineerd om een beeld te maken met een resolutie die gelijkwaardig is aan die van een radiotelescoop met een schoteldiameter van honderden kilometers. Het nadeel van deze combinatie van antennes is de benodigde rekenkracht voor het maken van dergelijke wetenschappelijke afbeeldingen.

De LOFAR (LOW-Frequency-ARray) radiotelescoop is een Nederlandse laagfrequente telescoop die bestaat uit meer dan 5000 antennes in Nederland, met duizenden meer gegroepeerd in stations in heel Europa. LOFAR observeert de radiohemel met de laagste frequenties zichtbaar vanaf de aarde: van 10 MHz tot 240 MHz. Het is ontworpen om gegevens te verzamelen voor meerdere wetenschappelijke toepassingen. Wetenschappers kunnen LOFAR-gegevens gebruiken om fenomenen te bestuderen zoals supernovaresten en pulsars in ons sterrenstelsel, magnetische velden van het zonnestelsel, samenvoegende melkwegclusters, zwarte gaten in het centrum van het verre sterrenstelsels en zelfs het tijdperk waarin de eerste sterren in ons universum werden gevormd. Om zulke uiteenlopende wetenschappelijke gevallen te dienen, worden LOFAR-gegevens opgeslagen met hoge tijd- en frequentie-resoluties. Dit leidt tot aanzienlijke gegevensgroottes. Elke observatie van 8 uur kan gemakkelijk acht harde schijven van 2 TB vullen. Het maken van een afbeelding op basis van deze gegevens is net aan mogelijk op iemands PC en een survey van de gehele hemel bestaat uit duizenden observaties die niet op één computer of zelfs een klein cluster van computers kunnen worden verwerkt. Desondanks kan elke verwerkingscyclus meerdere dagen duren, een vertraging die niet acceptabel is voor projecten die binnen vijf jaar duizenden datasets produceren.

Om duizenden waarnemingen van meerdere petabytes te kunnen verwerken, maken we gebruik van supercomputers voor parallelle verwerking van radioastronomiegegevens. Parallelliseren houdt in dat al onze gegevens kunnen worden opgesplitst in veel verschillende stukken, die elk onafhankelijk kunnen worden verwerkt. Bij een voldoende aantal computers daalt de verwerkingstijd met een factor 10 of meer. Het versnellen van de eerste paar verwerkingsstappen levert een tweede voordeel op: de eerste stappen reduceren de grootte van gegevens met een hoge resolutie tot 64 keer. Opeens past elke observatie gemakkelijk op uw desktop, laptop, en zelfs op een micro-SD kaart! Ze kunnen ook in seconden in plaats van uren tussen universiteiten worden uitgewisseld.

Dit werk is erop gericht om wetenschappers in staat te stellen van wetenschappers om eenvoudig en snel LOFAR-gegevens te verwerken, waardoor het gemakkelijker wordt om grote gegevenssets te gebruiken om wetenschappelijke studies uit te voeren. We maken gebruik van open source software, een verwerkingsinfrastructuur met een hoge verwerkingscapaciteit en de eigenschappen van gegevensverwerking via radioastronomie om grote wetenschappelijke onderzoeken met LOFAR mogelijk te maken. Ons werk is echter niet alleen nuttig voor wetenschappers. We bouwen en presenteren tools waarmee we de softwareprestaties van complexe wetenschappelijke pijpleidingen kunnen bestuderen en gebruiken deze tools om efficiëntere verwerkingsstrategieën aan te bevelen. Onze resultaten zijn ook van toepassing op toekomstige verwerkingsprojecten, zowel voor LOFAR-gegevens als voor de volgende generatie radiotelescopen. Uiteindelijk kunnen de lessen die zijn geleerd van de uitgebreide LOFAR-observaties worden gebruikt voor toekomstige astronomische projecten met grote hoeveelheden gegevens.

In Hoofdstuk Één geven we een overzicht van de geschiedenis van wetenschappelijk computergebruik, radio-interferometrie en de uitdagingen bij de verwerking van LOFAR-gegevens. In Hoofdstuk Twee presenteren we een platform voor grootschalige gedistribueerde LOFAR-verwerking. Dit platform is gebouwd voor geavanceerde LOFAR-gebruikers die hun verwerking willen paralleliseren op een gedeeld systeem, verdeeld over meerdere computationele toestellen en clusters. We introduceren de gebruikte softwarepakketten en hun interacties en bespreken de noodzaak van een dergelijk platform voor huidige en toekomstige grootschalige astronomische studies.

In Hoofdstuk Drie introduceren we een raamwerk voor het starten, volgen en paralleliseren van verwerkingsopdrachten voor de LOFAR-telescoop. Het is de eerste keer dat LOFAR-gegevens in bulk werden verwerkt op een High Throughput Infrastructuur. We laten zien dat dit computerparadigma kan worden toegepast op

de eerste stappen van een enkele verwerkingspijplijn. We geven een suggestie hoe andere pijpleidingen en telescopen dit kader kunnen gebruiken om hun gegevens te verwerken. Ten slotte laten we een aanzienlijke versnelling van de gegevensverwerking zien in vergelijking met de gegevensverwerking in Leiden: tot 35 keer sneller.

In Hoofdstuk Vier gaan we de uitdaging aan om prestatiegegevens te verzamelen van gedistribueerde runs van een complexe pijplijn. Voortbouwend op ons werk met betrekking tot parallelle verwerkingspijplijnen, bouwen we software om de prestaties van elke verwerkingsstap te volgen. De gegevens die we bijhouden, worden verzameld op een centrale server en kunnen in realtime worden geanalyseerd of worden bestudeerd na de verwerkingsrun. We voeren tests uit op vier verschillende systemen en constateren dat de compilatiemethode de prestaties van de verwerking niet verslechtert. We krijgen ook inzicht in de prestaties de lage prestaties van onze langzaamste stappen. Onze bewakingssoftware biedt een rijke dataset voor wetenschappelijke softwareontwikkelaars om inzicht te krijgen in de realtime prestaties van de gegevensreductiepijplijnen.

In Hoofdstuk Vijf introduceren we de eerste workflow-orkestratiesoftware voor complexe LOFAR-pijpleidingen. Deze software is het best geschikt voor de automatische verwerking van LOFAR-gegevens geproduceerd door grote, langlopende projecten. Deze software combineert al ons eerdere werk en maakt het eenvoudiger om grote verwerkingspijplijnen in te zetten op een grootschalige gedistribueerde infrastructuur. Het doel is om complexe LOFAR-pijpleidingen eenvoudig te visualiseren en te automatiseren. De verwerking van de 3000+ LOFAR twee-meter Sky Survey is geautomatiseerd met onze software, waaruit blijkt dat het complexe verwerking, parallelisatie, gegevensarchivering en externe databasetoegang aankan. Onze tool is gebouwd om substantiële, petabyte-grootte, LOFAR-projecten mogelijk te maken.

In Hoofdstuk Zes demonstreren we een methode voor het maken van een compleet schaalbaarheidsmodel voor complexe pijpleidingen. Met de LOFAR-prefactor-pijplijn omvatten we meer dan een orde van grootte in verwerkingsparameters om te begrijpen hoe onze software presteert met toenemende gegevensgroottes. De geleerde lessen helpen onze huidige gegevensverwerking te optimaliseren, de tijd te voorspellen die toekomstige taken in zullen beslag nemen en inzicht te krijgen in de prestaties van grote gegevenssets.

Hoofdstuk Zeven beschrijft een methode om ervoor te zorgen dat toekomstige softwarebeelden vóór de implementatie worden getest tegen productiepijplijnen. We automatiseren dit met behulp van het workflow-orkestratiesysteem

beschreven in Hoofdstuk vijf. In het laatste hoofdstuk beantwoorden we de onderzoeksvragen in onze inleiding en bespreken we de beperkingen van ons werk. Eindelijk, sluiten we af met een bespreking van toekomstige toepassingen van onze software.



## English Summary

---

When we look up at the night sky, we observe the Universe through visible light. We can see stars, planets, nebulae, and maybe even other galaxies. The earliest astronomers studied the Universe through their eyes, eventually aided by mirrors and lenses in the form of telescopes and binoculars. In the 18th and 19th centuries, physicists have unraveled the breadth of the electromagnetic spectrum. Aside from visible light, astronomers were now able to study the Universe using the infrared, ultra-violet, microwave, and radio parts of the spectrum. Light from each part of the spectrum carries different information with it and can be used to study specific physical phenomena. The drawback to this breadth of information is that every part of the spectrum needs its own dedicated detectors, and often, entire telescopes. Nevertheless, each part of the spectrum reveals a part of the Universe hidden to us in all other wavelengths.

Radio Astronomy was born in the 1930s with Karl Jansky's experiments with a directed 30-meter radio antenna. With it, he was able to detect thunderstorms, the Sun's magnetosphere, but also a strange unknown source at the center of our galaxies. These experiments proved that low-frequency radio waves could be used to study the distant Universe. In the 1940s, radar receivers were developed as a matter of national security during the Second World War. They are considered the main reason for the victory in the Battle of Britain over the Luftwaffe. After the conclusion of the war, some of the radar hardware was turned to the skies.

While the new antennas were much more sensitive to radio frequencies, radio telescopes have a fundamental limitation in terms of resolution. Because light behaves as a wave, especially at long wavelengths, it is bound by the diffraction limit. The diffraction limit links the smallest object that could be spatially resolved by a telescope with the wavelength of light and the telescope's diameter. If we take a 100-m radio telescope, it will have the resolving accuracy equivalent to a visible telescope with a mirror with a 0.5 mm diameter. This shouldn't be too shocking: radio waves

are more than 10 million times longer than the light we perceive, so naturally, our telescopes need to be 10 million times larger.

To match the resolution of the Hubble space telescope, the mirror of a radio telescope at 150 MHz would need a diameter of  $\sim 1000\text{km}$ . While engineers are still working on making radio dishes of this size, astronomers have turned to computers to increase the angular resolution of radio telescopes. Using the wave properties of light, radio astronomers can synchronize the timing of the data observed at multiple antennas, separated by tens, hundreds, or thousands of kilometers. With significant processing, data from these distant antennas can be combined together to make an image with a resolution equivalent to a radio telescope with a dish diameter of hundreds of kilometers. The downside of this combination of antennas is the computational requirements of making such scientific images.

The LOFAR (LOw-Frequency-ARray) radio telescope is a Dutch low-frequency telescope that consists of more than 5000 antennas in the Netherlands, with thousands more grouped in stations across Europe. LOFAR observes the radio sky at the lowest frequencies visible from Earth: from 10 MHz to 240 MHz. It is designed to collect data for multiple science cases. Scientists can use LOFAR data to study phenomena such as supernova remnants and pulsars in our galaxy, solar system magnetic fields, merging galaxy clusters, black holes in the centre of distant galaxies, and even the epoch when the first stars in our Universe formed. To serve such diverse science cases, LOFAR data is stored at high time and frequency resolutions. This leads to considerable data sizes. Each 8-hour observation can easily fill eight 2-TB hard drives. Creating an image from this data is just possible on one's personal computer and large all-sky surveys consist of thousands of observations which cannot be processed on a single computer or even a small cluster of computers. Even so, each run can take several days, a latency not feasible for projects producing several thousand data sets within a five-year project.

To make it possible to process thousands of multi-petabyte observations, we take advantage of supercomputers for parallel processing of radio astronomy data. Data parallelism means that all of our data can be split into many different pieces, each of which can be processed independently. With a sufficient number of computers, the processing time drops by a factor of 10 or more. Accelerating the first few processing steps delivers a second advantage: The initial stages take high-resolution data and average it down by a factor of up to 64. Suddenly each observation can comfortably sit on your desktop, laptop, and even on a micro-SD card! They can also be transported between universities in seconds rather than hours.

This work is focused on how to enable scientists to easily and quickly process LOFAR data, making it easier to use large data sets to conduct scientific studies. We use open source libraries, a high throughput processing infrastructure, and the properties of radio astronomy data processing to make large scientific surveys with LOFAR possible. Our work is not only useful for scientists, though. We build and present tools that enables us to study the software performance of complex scientific pipelines, and use these tools to recommend more efficient processing strategies. Our results are also applicable to future processing efforts, both for LOFAR data and for upcoming radio telescopes. Ultimately, the lessons learned from the extensive LOFAR surveys can be used for future astronomical projects tasked with large data sizes.

In Chapter One, we give an overview of the history of scientific computing, radio interferometry, and the processing challenges for LOFAR data. In Chapter Two, we present a platform for large scale distributed LOFAR processing. This platform is built for advanced LOFAR users who wish to parallelize their processing on a shared system distributed across multiple computational nodes and clusters. We present the software packages used and their interactions and discuss the necessity of such a platform for current and future large scale astronomical studies.

In Chapter Three, we introduce a framework for launching, tracking, and parallelizing processing jobs for the LOFAR telescope. Our results represent the first time that LOFAR data were processed in bulk, on a High Throughput Infrastructure. We show the applicability of this computing paradigm on the initial steps of a single processing pipeline. We suggest how other pipelines and telescopes can use this framework to process their data. Finally, we show a significant speed-up in data processing compared to data processing in Leiden: up to 35 times faster.

In Chapter Four, we tackle the challenge of collecting performance data from distributed runs of a complex pipeline. Building on our work parallelizing processing pipelines, we build software to track the performance of each processing step. The data that we track is collected at a central server and can be analysed in real-time or studied after the processing run. We run tests on four different systems and find that the software compilation method doesn't degrade the run time performance of the processing. We also gain insights into low-level performance for our slowest steps. Our monitoring software provides a rich data set for scientific software developers to gain insights into the real time performance of the data reduction pipelines.

In Chapter Five, we introduce the first workflow orchestration software for

complex LOFAR pipelines. This software is suited best for the automatic processing of LOFAR data produced by large, long-running projects. This software combines all of our previous work and makes it easier to deploy large processing pipelines on a large scale distributed infrastructure. We have built it to visualise and automate complex LOFAR pipelines easily. The processing of the 3000+ LOFAR Two-Metre Sky Survey is automated with our software, showing its ability to handle complex processing, parallelization, data archival, and remote database access. Our tool is built to make substantial, petabyte-scale, LOFAR projects feasible.

In Chapter Six, we demonstrate a method for creating a complete scalability model for complex pipelines. Using the LOFAR prefactor pipeline, we span more than an order of magnitude in processing parameters to understand how our software performs with increasing data sizes. The lessons learned will help optimize our current data processing, predict the time taken by future jobs, and understand the performance of our processing for large data sets.

Chapter Seven describes a method to ensure future software images are tested against production pipelines before deployment. We automate this using the workflow orchestration system described in Chapter five. In the final chapter, we answer the research questions posed in our introduction and discuss the limitations of our work and conclude with a discussion of future applications of our software.

### 8.5 Journal Articles

**A.P. Mechev**, A. Plaat, J.B.R. Oonk, H.T. Intema, and H.J.A. Röttgering. “Pipeline Collector: Gathering performance data for distributed astronomical pipelines”. In: *Astronomy and Computing* 24 (2018), pp. 117–128. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2018.06.005>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133718300490>

**A.P. Mechev**, T.W. Shimwell, A. Plaat, H.T. Intema, A.L. Varbanescu, and H.J.A. Röttgering. “Scalability model for the LOFAR direction independent pipeline”. In: *Astronomy and Computing* 28 (2019), p. 100293. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2019.100293>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133719300290>

**Alexandar P Mechev**, Aske Plaat, Huib Intema, and Huub Rottgering. “Automated testing and quality control of LOFAR scientific pipelines with AGLOW”. in: *Astronomy and Computing* (2019 in review)

JBR Oonk, **AP Mechev**, N Danezi, F Sweijen, T Shimwell, C Schrijvers, A Drabent, and K Emig. “Radio astronomical reduction on distributed and shared processing infrastructures: a platform for LOFAR”. in: *Astronomy and Computing* (2019 in prep)

T. W. Shimwell, Tasse, C., Hardcastle, M. J., **Mechev, A. P.**, Williams, W. L., Best, P. N., Röttgering, H. J. A., Callingham, J. R., Dijkema, T. J., de Gasperin, F., Hoang, D. N., Hugo, B., Mirmont, M., Oonk, J. B. R., Prandoni, I., Rafferty, D., Sabater, J., Smirnov, O., van Weeren, R. J., and White, G. J. et al. “The LOFAR Two-metre Sky Survey - II. First data release”. In: *A&A* 622 (2019), A1. DOI: 10.1051/0004-6361/201833559. URL: <https://doi.org/10.1051/0004-6361/201833559>

K. L. Emig, P. Salas, F. de Gasperin, J. B. R. Oonk, M. C. Toribio, **A. P. Mechev**, H. J. A. Röttgering, and A. G. G. M. Tielens. “Searching for the largest bound atoms in space”. In: *A&A* (2019)

## 8.6 Conference Proceedings

Y. G. Grange, R. Lakhoo, M. Petschow, C. Wu, B. Veenboer, I. Emsley, T. J. Dijkema, **A. P. Mechev**, and G. Mariani. “Characterising radio telescope software with the Workload Characterisation Framework”. In: *arXiv e-prints*, arXiv:1612.00456 (Dec. 2016), arXiv:1612.00456. arXiv: 1612.00456 [astro-ph.IM]

**A.P. Mechev**, J.B.R. Oonk, A. Danezi, T.W. Shimwell, C. Schrijvers, H.T. Intema, A. Plaat, and H.J.A. Röttgering. “An Automated Scalable Framework for Distributing Radio Astronomy Processing Across Clusters and Clouds”. In: *Proceedings of the International Symposium on Grids and Clouds (ISGC) 2017, held 5-10 March, 2017 at Academia Sinica, Taipei, Taiwan (ISGC2017)*. Online at <https://pos.sissa.it/cgi-bin/reader/conf.cgi?confid=293>, id.2. Mar. 2017, p. 2. arXiv: 1712.00312 [astro-ph.IM]

**A.P. Mechev**, J.B.R. Oonk, T.W. Shimwell, A. Plaat, H.T. Intema, and H.J.A. Röttgering. “Fast and Reproducible LOFAR Workflows with AGLOW”. in: *2018 IEEE 14th International Conference on e-Science (e-Science)*. Oct. 2018, arXiv:1808.10735. doi: 10.1109/eScience.2018.00029. arXiv: 1808.10735 [astro-ph.IM]

**A.P. Mechev**, J.B.R. Oonk, A. Plaat, N. Danezi, and T.W. Shimwell. “Building LOFAR as a Service”. In: *Astronomical Data Analysis Software and Systems XXVIII*, p. 677

**AGLOW** AGLOW is a combination of Apache Airflow with GRID\_LRT. This integration allows LOFAR users to build and launch massively parallel workflows. 84

**CouchDB** CouchDB is a document-based eventually consistent database, that we use to store processing information for distributed jobs. Each CouchDB document corresponds to a single distributed job, and contains a full description of the job required to run on a worker node. As jobs run, they update their status in the CouchDB document, which can be accessed by users through their browsers or a Python client. 24, 45, 91

**CVMFS** The CERN VM Filesystem is a virtually mounted filesystem that is used to distribute software on multiple clusters, cluster nodes and individual machines. CVMFS allows an institute to host a portable installation of their software, which is distributed and cached by other CVMFS clients. The software is cached locally on the worker nodes as a FileSystem in Userspace (FUSE) module . 26, 54, 67, 109

**dCache** dCache is a system for storing and retrieving large amounts of data, distributed across heterogeneous servers. dCache provides a common virtual filesystem, while also allows data to be located on varied storage devices including SSDs, spinning disks and magnetic tape . 152

**Grid** Grid computing refers to massively parallel distributed computing introduced in the '90s to tackle the processing challenges processing data from the Large Hadron Collider. A computational grid is a set of compute nodes connected with a high throughput connection, common job scheduler and shared, distributed storage. The computational and storage resources in a Grid are federated, and users are provided a share of those resources by a managing authority . 3, 18, 45, 61, 84, 104, 149

**GRID\_LRT** GRID\_LRT is the GRID LOFAR Reduction Tools package. This software consists of a set of tools to easily create and launch processing jobs on a distributed infrastructure. It includes tools to manage LOFAR data stored on the grid filesystem. These tools make it possible to quickly integrate processing scripts with a high throughput environment, accelerating bottleneck steps in LOFAR processing . 18

**HBA** The High Band Array is an array of LOFAR antennas sensitive to 110-240MHz. Each lofar station in the Netherlands has 48 HBA elements, with core stations having two separate sub-arrays of 24; while international stations have 96 HBA elements. The naming scheme of these antennas is as such: LLNNHBA<sub>S</sub>, where LL is the location (CS for core stations, RS for remote stations, and the ISO 3166-1 2-letter country code for the international stations); NNN is the station number, starting at 000, HBA/LBA denotes whether it's a HBA or LBA antenna and S refers to the core station sub-array. 6, 19

**HPC** High Performance Computing, as opposed to High Throughput Computing (HTC), refers to computation on one or multiple machines with many, fast CPUs; large quantities of RAM and fast disks. Often HPC jobs are done on clusters where multiple of these machines are connected with a fast network used to pass messages and to synchronize workloads. 20

**HTC** High Throughput Computing, as opposed to High Performance Computing (HPC), focuses on minimizing the time taken processing large amounts of data by using techniques in streaming, parallelizing and distributing many small jobs on a cluster . 20, 60, 84, 135

**LOFAR** The LOw Frequency ARray: A large, low-frequency aperture synthesis radio telescope. 6, 17, 44, 60, 84, 102, 134

**LoTSS** The LOFAR Two-Meter Sky Survey is a whole-sky study of the low-frequency radio sky at 120-168MHz. LoTSS is composed of a broad, Tier 1, survey of the entire sky, as well as deeper tiers targeted at specific fields of interest . 7, 18, 43, 44, 105, 136, 149

**srnm** Storage Resource Manager, a system to co-ordinate data storage. This system defines the protocols for referencing data on a distributed system, accessing metadata and changing data locality (e.g., moving from tape to disk). . 20, 50, 92

**Subband** Broadband LOFAR observations are stored in separate 'Subbands', splitting the frequency range into several individual files, before storing at the Long Term Archive. Depending on the observation mode, one observation can have 230-480 Subbands. This splitting makes it easier for users to request, download and process a fraction of observation's entire bandwidth. 8, 31, 47, 85, 104

**visibilities** Radio Astronomers use the term 'UV plane' or 'visibilities' interchangeably to refer to the Fourier Transform of the final image. Each baseline of an aperture synthesis telescope corresponds to one measurement in this space. The letters U and V refer to the two orthogonal components of a baseline with respect to an observation's phase center. The  $u$  and  $v$  vectors are defined in a plane orthogonal to the direction towards the phase center, and are typically in units of wavelength. To obtain an image, the UV data needs to be 'cleaned' by iteratively removing the point spread function of the telescope. 5, 20, 107

**VOMS** The Virtual Organization Membership Service is a service that manages the access to data and computational resources provided to each Grid user. It handles authentication and authorisation of job launching and access to data on the grid filesystem. More information can be found at <https://italiangrid.github.io/voms/>. 24



## Bibliography

---

- [1] John von Neumann. *First Draft of a Report on the EDVAC*. Tech. rep. 1945.
- [2] William W. Wood. *Early history of computer simulations in statistical mechanics*. Tech. rep. Los Alamos National Lab., NM (USA); Carroll Coll., Helena, MT (USA), 1985.
- [3] Edward N Lorenz. “Empirical orthogonal functions and statistical weather prediction”. In: (1956).
- [4] Francis H Harlow and J Eddie Welch. “Numerical calculation of time-dependent viscous incompressible flow of fluid with free surface”. In: *The physics of fluids* 8.12 (1965), pp. 2182–2189.
- [5] Frances E. Allen. “The history of language processor technology in IBM”. In: *IBM Journal of Research and Development* 25.5 (1981), pp. 535–548.
- [6] Christopher Bingham, M Godfrey, and J Tukey. “Modern techniques of power spectrum estimation”. In: *IEEE Transactions on audio and electroacoustics* 15.2 (1967), pp. 56–66.
- [7] Jack Dongarra and Francis Sullivan. “Guest Editors’ Introduction: The Top 10 Algorithms”. In: *Computing in Science & Engineering* 2.1 (2000), pp. 22–23. doi: 10.1109/MCISE.2000.814652. eprint: <https://aip.scitation.org/doi/pdf/10.1109/MCISE.2000.814652>. URL: <https://aip.scitation.org/doi/abs/10.1109/MCISE.2000.814652>.
- [8] Gerard Tel. *Introduction to Distributed Algorithms*. 2nd ed. Cambridge University Press, 2000. doi: 10.1017/CB09781139168724.
- [9] E. Fomalont. “Astronomical Image Processing System / AIPS”. In: *National Radio Astronomy Observatory Newsletter* 3 (1981), p. 3.
- [10] Doug Tody. “The IRAF data reduction and analysis system”. In: *Instrumentation in astronomy VI*. Vol. 627. International Society for Optics and Photonics. 1986, pp. 733–748.

- [11] Donald Carson Wells and Eric W Greisen. “FITS-a flexible image transport system”. In: *Image Processing in Astronomy*. 1979, p. 445.
- [12] Shang-Wen Cheng et al. “Software architecture-based adaptation for grid computing”. In: (1989).
- [13] BG Clark. “An efficient implementation of the algorithm ‘CLEAN’”. In: *Astronomy and Astrophysics* 89 (1980), p. 377.
- [14] Anthony Richard Thompson, James M Moran, George Warner Swenson, et al. *Interferometry and synthesis in radio astronomy*. Wiley New York et al., 1986.
- [15] Wilbur Norman Christiansen and Jan A Högbom. *Radiotelescopes*. CUP Archive, 1987.
- [16] F. de Gasperin et al. “Systematic effects in LOFAR data: A unified calibration strategy”. In: *Astronomy and Astrophysics* 622, A5 (Feb. 2019), A5. doi: 10.1051/0004-6361/201833867. arXiv: 1811.07954 [astro-ph.IM].
- [17] WL Williams et al. “LOFAR 150-MHz observations of the Boötes field: catalogue and source counts”. In: *Monthly Notices of the Royal Astronomical Society* 460.3 (2016), pp. 2385–2412.
- [18] R. J. van Weeren et al. “LOFAR Facet Calibration”. In: *ApJS* 223.1, 2 (Mar. 2016), p. 2. doi: 10.3847/0067-0049/223/1/2. arXiv: 1601.05422 [astro-ph.IM].
- [19] C. Tasse et al. “Faceting for direction-dependent spectral deconvolution”. In: *A&A* 611, A87 (Apr. 2018), A87. doi: 10.1051/0004-6361/201731474. arXiv: 1712.02078 [astro-ph.IM].
- [20] M. P. van Haarlem et al. “LOFAR: The LOw-Frequency ARray”. In: *A&A* 556, A2 (Aug. 2013), A2. doi: 10.1051/0004-6361/201220873. arXiv: 1305.3550 [astro-ph.IM].
- [21] *LOFAR Stations: Description and Layout*. URL: <https://old.astron.nl/radio-observatory/astronomers/%20users/technical-information/lofar-stations/lofar-stations-description->.
- [22] ASTRON. *LOFAR Brochure*. Apr. 2019. URL: [https://www.astron.nl/sites/default/files/shared/LOFAR\\_Brochure\\_ENG\\_April2019.pdf](https://www.astron.nl/sites/default/files/shared/LOFAR_Brochure_ENG_April2019.pdf).
- [23] AH Patil et al. “Upper limits on the 21 cm epoch of reionization power spectrum from one night with LOFAR”. In: *The Astrophysical Journal* 838.1 (2017), p. 65.
- [24] T. W. Shimwell et al. “The LOFAR Two-metre Sky Survey - II. First data release”. In: *arXiv e-prints*, arXiv:1811.07926 (Nov. 2018), arXiv:1811.07926. arXiv: 1811.07926 [astro-ph.GA].

- [25] TW Shimwell et al. “The LOFAR Two-metre Sky Survey-I. Survey description and preliminary data release”. In: *Astronomy & Astrophysics* 598 (2017), A104.
- [26] Heald, G. H. et al. “The LOFAR Multifrequency Snapshot Sky Survey (MSSS) - I. Survey description and first results”. In: *A&A* 582 (2015), A123. doi: 10 . 1051 / 0004 - 6361 / 201425210. URL: [https : / / doi . org / 10 . 1051 / 0004 - 6361 / 201425210](https://doi.org/10.1051/0004-6361/201425210).
- [27] J. B. R. Oonk et al. “Carbon and hydrogen radio recombination lines from the cold clouds towards Cassiopeia A”. In: *MNRAS* 465.1 (Feb. 2017), pp. 1066–1088. doi: 10.1093/mnras/stw2818. arXiv: 1609.06857 [astro-ph.GA].
- [28] K. L. Emig et al. “The first detection of radio recombination lines at cosmological distances”. In: *A&A* 622, A7 (Feb. 2019), A7. doi: 10 . 1051 / 0004 - 6361 / 201834052. arXiv: 1811.08104 [astro-ph.GA].
- [29] M. Arias et al. “Low-frequency radio absorption in Cassiopeia A”. In: *A&A* 612, A110 (Apr. 2018), A110. doi: 10 . 1051 / 0004 - 6361 / 201732411. arXiv: 1801 . 04887 [astro-ph.HE].
- [30] P. Salas et al. “Mapping low-frequency carbon radio recombination lines towards Cassiopeia A at 340, 148, 54, and 43 MHz”. In: *MNRAS* 475.2 (Apr. 2018), pp. 2496–2511. doi: 10 . 1093 / mnras / stx3340. arXiv: 1801 . 05298 [astro-ph.GA].
- [31] M. E. Bell et al. “An automated archival Very Large Array transients survey”. In: *MNRAS* 415 (July 2011), pp. 2–10. doi: 10.1111/j.1365-2966.2011.18631.x. arXiv: 1103.0511 [astro-ph.HE].
- [32] B. W. Stappers et al. “Observing pulsars and fast transients with LOFAR”. In: *A&A* 530, A80 (June 2011), A80. doi: 10 . 1051 / 0004 - 6361 / 201116681. arXiv: 1104 . 1577 [astro-ph.IM].
- [33] G.N.J. van Diepen. “Casacore Table Data System and its use in the MeasurementSet”. In: *Astronomy and Computing* 12 (2015), pp. 174–180. ISSN: 2213-1337. doi: [https : / / doi . org / 10 . 1016 / j . ascom . 2015 . 06 . 002](https://doi.org/10.1016/j.ascom.2015.06.002). URL: [http : / / www . sciencedirect . com / science / article / pii / S2213133715000621](http://www.sciencedirect.com/science/article/pii/S2213133715000621).
- [34] M. J. Hardcastle et al. “LOFAR/H-ATLAS: a deep low-frequency survey of the Herschel-ATLAS North Galactic Pole field”. In: *MNRAS* 462.2 (Oct. 2016), pp. 1910–1936. doi: 10 . 1093 / mnras / stw1763. arXiv: 1606 . 09437 [astro-ph.GA].

- [35] **A.P. Mechev** et al. “An Automated Scalable Framework for Distributing Radio Astronomy Processing Across Clusters and Clouds”. In: *Proceedings of the International Symposium on Grids and Clouds (ISGC) 2017, held 5-10 March, 2017 at Academia Sinica, Taipei, Taiwan (ISGC2017)*. Online at <https://pos.sissa.it/cgi-bin/reader/conf.cgi?confid=293,id.2>. Mar. 2017, p. 2. arXiv: 1712.00312 [astro-ph.IM].
- [36] **A.P. Mechev** et al. “Pipeline Collector: Gathering performance data for distributed astronomical pipelines”. In: *Astronomy and Computing* 24 (2018), pp. 117–128. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2018.06.005>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133718300490>.
- [37] **A.P. Mechev** et al. “Fast and Reproducible LOFAR Workflows with AGLOW”. In: *2018 IEEE 14th International Conference on e-Science (e-Science)*. Oct. 2018, arXiv:1808.10735. DOI: 10.1109/eScience.2018.00029. arXiv: 1808.10735 [astro-ph.IM].
- [38] **A.P. Mechev** et al. “Scalability model for the LOFAR direction independent pipeline”. In: *Astronomy and Computing* 28 (2019), p. 100293. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2019.100293>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133719300290>.
- [39] C.L. Carilli and S. Rawlings. “Motivation, key science projects, standards and assumptions”. In: *New Astronomy Reviews* 48.11 (2004). Science with the Square Kilometre Array, pp. 979–984. ISSN: 1387-6473. DOI: <https://doi.org/10.1016/j.newar.2004.09.001>. URL: <http://www.sciencedirect.com/science/article/pii/S1387647304000880>.
- [40] Andreas Horneffer et al. *apmechev/prefactor: LOTSS Data Release 1*. Nov. 2018. DOI: 10.5281/zenodo.1487962. URL: <https://doi.org/10.5281/zenodo.1487962>.
- [41] T. W. Shimwell et al. “The LOFAR Two-metre Sky Survey. I. Survey description and preliminary data release”. In: *A&A* 598, A104 (Feb. 2017), A104. DOI: 10.1051/0004-6361/201629313. arXiv: 1611.02700 [astro-ph.IM].
- [42] A. Wilber et al. “LOFAR discovery of an ultra-steep radio halo and giant head-tail radio galaxy in Abell 1132”. In: *MNRAS* 473.3 (Jan. 2018), pp. 3536–3546. DOI: 10.1093/mnras/stx2568. arXiv: 1708.08928 [astro-ph.GA].
- [43] H. J. A. Rottgering et al. “LOFAR - Opening up a new window on the Universe”. In: *arXiv e-prints*, astro-ph/0610596 (Oct. 2006), astro-ph/0610596. arXiv: astro-ph/0610596 [astro-ph].

- [44] Offringa, A. R. et al. “The LOFAR radio environment”. In: *A&A* 549 (2013), A11. DOI: 10.1051/0004-6361/201220293. URL: <https://doi.org/10.1051/0004-6361/201220293>.
- [45] Tammo Jan Dijkema. “RFI Flagging, Demixing and Visibilities Compression”. In: *Astrophysics and Space Science Library*. Vol. 426. Astrophysics and Space Science Library. Jan. 2018, p. 55. DOI: 10.1007/978-3-319-23434-2\_4.
- [46] Jamie Shiers. “The Worldwide LHC Computing Grid (worldwide LCG)”. In: *Computer Physics Communications* 177.1 (2007). Proceedings of the Conference on Computational Physics 2006, pp. 219–223. ISSN: 0010-4655. DOI: <https://doi.org/10.1016/j.cpc.2007.02.021>. URL: <http://www.sciencedirect.com/science/article/pii/S001046550700077X>.
- [47] S. . van der Tol, B. D. Jeffs, and A. -J. . van der Veen. “Self-Calibration for the LOFAR Radio Astronomical Array”. In: *IEEE Transactions on Signal Processing* 55.9 (Sept. 2007), pp. 4497–4510. DOI: 10.1109/TSP.2007.896243.
- [48] Hanno Holties, Adriaan Renting, and Yan Grange. “The LOFAR long-term archive: e-infrastructure on petabyte scale”. In: *SPIE Astronomical Telescopes+ Instrumentation*. International Society for Optics and Photonics. 2012, pp. 845117–845117.
- [49] Domingos Barbosa et al. “Power Monitoring and Control for Large Scale projects: SKA, a case study”. In: *SPIE Astronomical Telescopes+ Instrumentation*. International Society for Optics and Photonics. 2016, pp. 99100L–99100L.
- [50] JBR Oonk et al. “Radio astronomical reduction on distributed and shared processing infrastructures: a platform for LOFAR”. In: *Astronomy and Computing* (2019 in prep).
- [51] Shayan Shams et al. “A Scalable Pipeline For Transcriptome Profiling Tasks With On-demand Computing Clouds”. In: *Parallel and Distributed Processing Symposium Workshops, 2016 IEEE International*. IEEE. 2016, pp. 443–452.
- [52] Anjani Ragotheraman et al. “Developing ethread pipeline using saga-pilot abstraction for large-scale structural bioinformatics”. In: *BioMed research international* 2014 (2014).
- [53] JD Van Horn et al. *Grid-Based Computing and the Future of Neuroscience Computation, Methods in Mind*. 2005.
- [54] William K Michener and Matthew B Jones. “Ecoinformatics: supporting ecology as a data-intensive science”. In: *Trends in ecology & evolution* 27.2 (2012), pp. 85–93.
- [55] Ewa Deelman et al. “Pegasus: A framework for mapping complex scientific workflows onto distributed systems”. In: *Scientific Programming* 13.3 (2005), pp. 219–237.

- [56] Yong Zhao et al. “Swift: Fast, reliable, loosely coupled parallel computation”. In: *Services, 2007 IEEE Congress on*. IEEE. 2007, pp. 199–206.
- [57] *PiCaS Overview - Grid Documentation v1.0*. [http://doc.grid.surfsara.nl/en/latest/Pages/Practices/picas/picas\\_overview.html](http://doc.grid.surfsara.nl/en/latest/Pages/Practices/picas/picas_overview.html). 2017.
- [58] C Aguado Sanchez et al. “CVMFS-a file system for the CernVM virtual appliance”. In: *Proceedings of XII Advanced Computing and Analysis Techniques in Physics Research*. Vol. 1. 2008, p. 52.
- [59] Jan Bot. *PiCaS: Python client using CouchDB as a token pool server*. <https://github.com/jjbot/picasclient>. 2017.
- [60] Chirs ANDERSON. “Apache couchdb: The definitive guide”. In: <http://couchdb.apache.org/index.htm> Acessado em 5.06 (2009), p. 2009.
- [61] Joris Borgdorff, Harsha Krishna, and Michael H Lees. “SIM-CITY: An e-Science framework for Urban Assisted Decision Support”. In: *Procedia Computer Science* 51 (2015), pp. 2327–2336.
- [62] Y Li. “Reliability of long heterogeneous slopes in 3d: Model performance and conditional simulation”. In: (2017).
- [63] Marco Clemencic and B Couturier. “A New Nightly Build System for LHCb”. In: *Journal of Physics: Conference Series*. Vol. 513. 5. IOP Publishing. 2014, p. 052007.
- [64] Tom SANTE. “Development of (graphical) web applications for the processing and interpretation of arrayCGH data”. In: (2010).
- [65] Grigory Rybkin. “ATLAS software packaging”. In: *Journal of Physics: Conference Series*. Vol. 396. 5. IOP Publishing. 2012, p. 052063.
- [66] GS Davies et al. “Software Management for the NOvAExperiment”. In: *Journal of Physics: Conference Series*. Vol. 664. 6. IOP Publishing. 2015, p. 062011.
- [67] WN Brouw. “Aperture synthesis”. In: *Image Processing Techniques in Astronomy*. Springer, 1975, pp. 301–307.
- [68] Cyril Tasse et al. “LOFAR calibration and wide-field imaging”. In: *Comptes Rendus Physique* 13.1 (2012), pp. 28–32.
- [69] RF Pizzo et al. “The Lofar Imaging Cookbook v2.0”. In: *internal ASTRON report* (2010).
- [70] *Gina specifications - GRID Documentation v1.0*. Available at [http://docs.surfsaralabs.nl/projects/grid/en/latest/Pages/Service/system\\_specifications/gina\\_specs.html](http://docs.surfsaralabs.nl/projects/grid/en/latest/Pages/Service/system_specifications/gina_specs.html). 2017.

- [71] *Lofar Software Stack*. [http://www.lofar.org/wiki/doku.php?id=public:software\\_stack\\_installation](http://www.lofar.org/wiki/doku.php?id=public:software_stack_installation). 2017.
- [72] *Softdrive on the GRID*. Available at [http://docs.surfsaralabs.nl/projects/grid/en/latest/Pages/Advanced/grid\\_software.html#softdrive](http://docs.surfsaralabs.nl/projects/grid/en/latest/Pages/Advanced/grid_software.html#softdrive).
- [73] Jason Maassen et al. *Xenon: Xenon 1.1.0*. Dec. 2015. doi: 10.5281/zenodo.35415. URL: <https://doi.org/10.5281/zenodo.35415>.
- [74] Software Foundation Apache. *TCollector: OpenTSDB documentation*. Available at [http://opentsdb.net/docs/build/html/user\\_guide/utilities/tcollector.html](http://opentsdb.net/docs/build/html/user_guide/utilities/tcollector.html). 2017.
- [75] P. Chris Broekema et al. “Cobalt: A GPU-based correlator and beamformer for LOFAR”. In: *Astronomy and Computing* 23 (2018), pp. 180–192. ISSN: 2213-1337. DOI: <https://doi.org/10.1016/j.ascom.2018.04.006>. URL: <http://www.sciencedirect.com/science/article/pii/S2213133717301439>.
- [76] Y Gupta et al. “The upgraded GMRT: opening new windows on the radio Universe”. In: *Current Science* 113.4 (2017), p. 707.
- [77] MP Van Haarlem et al. “LOFAR: The low-frequency array”. In: *Astronomy & astrophysics* 556 (2013), A2.
- [78] Colin J Lonsdale et al. “The munchison widefield array: Design overview”. In: *Proceedings of the IEEE* 97.8 (2009), pp. 1497–1506.
- [79] S J Tingay et al. “The Murchison widefield array: The square kilometre array precursor at low radio frequencies”. In: *Publications of the Astronomical Society of Australia* 30 (2013).
- [80] Justin L Jonas. “MeerKAT-The South African array with composite dishes and wide-band single pixel feeds”. In: *Proceedings of the IEEE* 97.8 (2009), pp. 1522–1530.
- [81] Chen Wu et al. “Optimising NGAS for the MWA Archive”. In: *Experimental Astronomy* 36.3 (2013), pp. 679–694.
- [82] David B Davidson. “Potential technological spin-offs from MeerKAT and the South African Square Kilometre Array bid”. In: *South African Journal of Science* 108.1-2 (2012), pp. 01–03.
- [83] OM Smirnov and Cyril Tasse. “Radio interferometric gain calibration as a complex optimization problem”. In: *Monthly Notices of the Royal Astronomical Society* 449.3 (2015), pp. 2668–2684.

- [84] J. B. R. Oonk et al. “Discovery of carbon radio recombination lines in absorption towards Cygnus A”. In: *MNRAS* 437 (Feb. 2014), pp. 3506–3515. doi: 10.1093/mnras/stt2158. arXiv: 1401.2876.
- [85] SURF. *Grid at SURFsara*. <https://www.surf.nl/en/services-and-products/grid/index.html>. 2018.
- [86] R Centeno et al. “The Helioseismic and Magnetic Imager (HMI) vector magnetic field pipeline: optimization of the spectral line inversion code”. In: *Solar Physics* 289.9 (2014), pp. 3531–3547.
- [87] Stephen Strother et al. “Optimizing the fMRI data-processing pipeline using prediction and reproducibility performance metrics: I. A preliminary group analysis”. In: *Neuroimage* 23 (2004), S196–S207.
- [88] S. Ott. “The Herschel Data Processing System — HIPE and Pipelines — Up and Running Since the Start of the Mission”. In: *Astronomical Data Analysis Software and Systems XIX*. Vol. 434. Dec. 2010, p. 139.
- [89] Jens-S Vöckler et al. “Kickstarting remote applications”. In: *2nd International Workshop on Grid Computing Environments*. 2006, pp. 1–8.
- [90] Matthew L Massie, Brent N Chun, and David E Culler. “The ganglia distributed monitoring system: design, implementation, and experience”. In: *Parallel Computing* 30.7 (2004), pp. 817–840.
- [91] Terrehon Bowden. *THE /proc FILESYSTEM v1.3*. <https://www.kernel.org/doc/Documentation/filesystems/proc.txt>. 2009.
- [92] Philip J Mucci et al. “PAPI: A portable interface to hardware performance counters”. In: *Proceedings of the department of defense HPCMP users group conference*. Vol. 710. 1999.
- [93] Brendan Gregg and Jim Mauro. *DTrace: Dynamic Tracing in Oracle Solaris, Mac OS X and FreeBSD*. Prentice Hall Professional, 2011.
- [94] Nicholas Nethercote and Julian Seward. “Valgrind: a framework for heavyweight dynamic binary instrumentation”. In: *ACM Sigplan notices*. Vol. 42. 6. ACM. 2007, pp. 89–100.
- [95] *LOFAR Imaging Cookbook*. Available at [http://www.astron.nl/sites/astron.nl/files/cms/lofar\\_imaging\\_cookbook\\_v19.pdf](http://www.astron.nl/sites/astron.nl/files/cms/lofar_imaging_cookbook_v19.pdf).
- [96] B Sigoure. *OpenTSDB scalable time series database (TSDB)*. 2012.
- [97] Dong H Ahn. *Measuring FLOPS using hardware performance counter technologies on LC systems*. Tech. rep. Lawrence Livermore National Laboratory (LLNL), Livermore, CA, 2008.

- [98] GM Loose. “LOFAR self-calibration using a blackboard software architecture”. In: *Astronomical Data Analysis Software and Systems XVII*. Vol. 394. 2008, p. 91.
- [99] Ger van Diepen and Tammo Jan Dijkema. *DPPP: Default Pre-Processing Pipeline*. Astrophysics Source Code Library. Apr. 2018. ascl: 1804.003.
- [100] Jakob Blomer et al. “Distributing LHC application software and conditions databases using the CernVM file system”. In: *Journal of Physics: Conference Series*. Vol. 331. 4. IOP Publishing. 2011, p. 042003.
- [101] Randy H. Katz and David A. Patterson. *Memory Hierarchy, CMPUT429/CMPE382 Winter 2001*. University of Calgary. Available at <https://webdocs.cs.ualberta.ca/~amaral/courses/429/webslides/Topic4-MemoryHierarchy/sld003.htm>. 2001.
- [102] Peter J Denning. “The working set model for program behavior”. In: *Communications of the ACM* 11.5 (1968), pp. 323–333.
- [103] Kim Hazelwood and James E Smith. “Exploring code cache eviction granularities in dynamic optimization systems”. In: *Code Generation and Optimization, 2004. CGO 2004. International Symposium on*. IEEE. 2004, pp. 89–99.
- [104] Kazushige Goto and Robert A Geijn. “Anatomy of high-performance matrix multiplication”. In: *ACM Transactions on Mathematical Software (TOMS)* 34.3 (2008), p. 12.
- [105] Stefano Salvini and Stefan J Wijnholds. “StEFCal-An Alternating Direction Implicit method for fast full polarization array calibration”. In: *General Assembly and Scientific Symposium (URSI GASS), 2014 XXXIth URSI*. IEEE. 2014, pp. 1–4.
- [106] Kevin Skadron et al. “Branch prediction, instruction-window size, and cache size: Performance tradeoffs and simulation techniques”. In: *IEEE Transactions on Computers* 48.11 (1999), pp. 1260–1281.
- [107] Tarush Jain and Tanmay Agrawal. “The haswell microarchitecture-4th generation processor”. In: *International Journal of Computer Science and Information Technologies* 4.3 (2013), pp. 477–480.
- [108] Subhradyuti Sarkar and Dean Tullsen. “Compiler techniques for reducing data cache miss rate on a multithreaded architecture”. In: *High Performance Embedded Architectures and Compilers* (2008), pp. 353–368.
- [109] Alvin R Lebeck et al. “A large, fast instruction window for tolerating cache misses”. In: *Computer Architecture, 2002. Proceedings. 29th Annual International Symposium on*. IEEE. 2002, pp. 59–70.

- [110] Shiliang Hu et al. "An approach for implementing efficient superscalar CISC processors". In: *High-Performance Computer Architecture, 2006. The Twelfth International Symposium on*. IEEE. 2006, pp. 41–52.
- [111] P Chris Broekema, Rob V van Nieuwpoort, and Henri E Bal. "The Square Kilometre Array science data processor. Preliminary compute platform design". In: *Journal of Instrumentation* 10.07 (2015), p. C07004.
- [112] A Patterson David and L Hennessy John. "Computer organization and design: the hardware/software interface". In: *San mateo, CA: Morgan Kaufmann Publishers* 1 (2005), p. 998.
- [113] Team Apache HBase. "Apache hbase reference guide". In: *Apache, version 2.0* (2015).
- [114] J Sabater et al. "Calibration of radio-astronomical data on the cloud. LOFAR, the pathway to SKA." In: *Highlights of Spanish Astrophysics VIII*. 2015, pp. 840–843.
- [115] Jeff Templon and Jan Bot. "The Dutch National e-Infrastructure". In: *International Symposium on Grids and Clouds (ISGC)*. Vol. 13. 18. 2016.
- [116] Erwin Laure et al. *Middleware for the next generation Grid infrastructure*. Tech. rep. CERN, 2004.
- [117] Jamie Shiers. "The worldwide LHC computing grid (worldwide LCG)". In: *Computer physics communications* 177.1-2 (2007), pp. 219–223.
- [118] Ilkay Altintas et al. "Kepler: an extensible system for design and execution of scientific workflows". In: *Scientific and Statistical Database Management, 2004. Proceedings. 16th International Conference on*. IEEE. 2004, pp. 423–424.
- [119] David Churches et al. "Programming scientific and distributed workflow with Triana services". In: *Concurrency and Computation: Practice and Experience* 18.10 (2006), pp. 1021–1037.
- [120] Ji Liu et al. "A survey of data-intensive scientific workflow management". In: *Journal of Grid Computing* 13.4 (2015), pp. 457–493.
- [121] Lin Wang. "Directed acyclic graph". In: *Encyclopedia of Systems Biology*. Springer, 2013, pp. 574–574.
- [122] David J Pearce and Paul HJ Kelly. "A dynamic topological sort algorithm for directed acyclic graphs". In: *Journal of Experimental Algorithmics (JEA)* 11 (2007), pp. 1–7.
- [123] A. B. Kahn. "Topological Sorting of Large Networks". In: *Commun. ACM* 5.11 (Nov. 1962), pp. 558–562. ISSN: 0001-0782. DOI: 10.1145/368996.369025. URL: <http://doi.acm.org/10.1145/368996.369025>.

- [124] Jianjun Zhou and Martin Müller. “Depth-first discovery algorithm for incremental topological sorting of directed acyclic graphs”. In: *Information Processing Letters* 88.4 (2003), pp. 195–200.
- [125] John Vivian et al. “Toil enables reproducible, open source, big biomedical data analyses”. In: *Nature biotechnology* 35.4 (2017), p. 314.
- [126] Paolo Di Tommaso et al. “Nextflow enables reproducible computational workflows”. In: *Nature biotechnology* 35.4 (2017), pp. 316–319.
- [127] W. Freudling et al. “Automated data reduction workflows for astronomy. The ESO Reflex environment”. In: *Astronomy & Astrophysics* 559, A96 (Nov. 2013), A96. doi: 10.1051/0004-6361/201322494. arXiv: 1311.5411 [astro-ph.IM].
- [128] Peter Amstutz et al. “Common Workflow Language, v1. 0”. In: (2016).
- [129] Michael Kotliar, Andrey Kartashov, and Artem Barski. “CWL-Airflow: a lightweight pipeline manager supporting Common Workflow Language”. In: *bioRxiv* (2018), p. 249243.
- [130] Patrick Fuhrmann and Volker Gölzow. “dCache, storage system for the future”. In: *European Conference on Parallel Processing*. Springer. 2006, pp. 1106–1113.
- [131] Linus Torvalds and Junio Hamano. “Git: Fast version control system”. In: *URL <http://git-scm.com>* (2010).
- [132] A. R. Offringa et al. “WSCLEAN: an implementation of a fast, generic wide-field imager for radio astronomy”. In: *Monthly Notices of the Royal Astronomical Society* 444 (Oct. 2014), pp. 606–619. doi: 10.1093/mnras/stu1368. arXiv: 1407.1943 [astro-ph.IM].
- [133] A. R. Offringa, J. J. van de Gronde, and J. B. T. M. Roerdink. “A morphological algorithm for improving radio-frequency interference detection”. In: *Astronomy & astrophysics* 539, A95 (Mar. 2012), A95. doi: 10.1051/0004-6361/201118497. arXiv: 1201.3364 [astro-ph.IM].
- [134] JP McMullin et al. “CASA architecture and applications”. In: *Astronomical data analysis software and systems XVI*. Vol. 376. 2007, p. 127.
- [135] Niruj Mohan and David Rafferty. “PyBDSF: Python Blob Detection and Source Finder”. In: *Astrophysics Source Code Library* (2015).
- [136] Cyril Tasse. “Nonlinear Kalman filters for calibration in radio interferometry”. In: *Astronomy & Astrophysics* 566 (2014), A127.
- [137] P. Salas et al. “LOFAR observations of decameter carbon radio recombination lines towards Cassiopeia A”. In: *Monthly Notices of the Royal Astronomical Society* 467 (May 2017), pp. 2274–2287. doi: 10.1093/mnras/stx239. arXiv: 1701.08802.

- [138] J.B.R. Oonk et al. "Radio astronomy on a distributed shared computing platform: The LOFAR case". In: (2018).
- [139] Peter E Dewdney et al. "The square kilometre array". In: *Proceedings of the IEEE* 97.8 (2009), pp. 1482–1496.
- [140] Jeff Templon and Jan Bot. "The Dutch National e-Infrastructure". To appear in Proceedings of Science edition of the International Symposium on Grids and Clouds (ISGC) 2016 13-18 March 2016, Academia Sinica, Taipei, Taiwan. Oct. 2016. URL: <https://doi.org/10.5281/zenodo.163537>.
- [141] H.A. Sanjay and Sathish Vadhiyar. "Performance modeling of parallel applications for grid scheduling". In: *Journal of Parallel and Distributed Computing* 68.8 (2008), pp. 1135–1145. ISSN: 0743-7315. DOI: <https://doi.org/10.1016/j.jpdc.2008.02.006>. URL: <http://www.sciencedirect.com/science/article/pii/S0743731508000464>.
- [142] Zhichen Xu, Xiaodong Zhang, and Lin Sun. "Semi-empirical Multiprocessor Performance Predictions". In: *Journal of Parallel and Distributed Computing* 39.1 (1996), pp. 14–28. ISSN: 0743-7315. DOI: <https://doi.org/10.1006/jpdc.1996.0151>. URL: <http://www.sciencedirect.com/science/article/pii/S0743731596901513>.
- [143] Bradley J Barnes et al. "A regression-based approach to scalability prediction". In: *Proceedings of the 22nd annual international conference on Supercomputing*. ACM. 2008, pp. 368–377.
- [144] Michael Kuperberg, Klaus Krogmann, and Ralf Reussner. "Performance prediction for black-box components using reengineered parametric behaviour models". In: *International Symposium on Component-Based Software Engineering*. Springer. 2008, pp. 48–63.
- [145] Carl Witt et al. "Predictive Performance Modeling for Distributed Computing using Black-Box Monitoring and Machine Learning". In: *CoRR* abs/1805.11877 (2018).
- [146] Leo T Yang, Xiaosong Ma, and Frank Mueller. "Cross-platform performance prediction of parallel applications using partial execution". In: *Proceedings of the 2005 ACM/IEEE conference on Supercomputing*. IEEE Computer Society. 2005, p. 40.
- [147] S. Kavulya et al. "An Analysis of Traces from a Production MapReduce Cluster". In: *2010 10th IEEE/ACM International Conference on Cluster, Cloud and Grid Computing*. May 2010, pp. 94–103. DOI: 10.1109/CCGRID.2010.112.

- [148] Laura Carrington, Allan Snavely, and Nicole Wolter. “A performance prediction framework for scientific applications”. In: *Future Generation Computer Systems* 22.3 (2006), pp. 336–346.
- [149] Alexandru Calotoiu et al. “Using automated performance modeling to find scalability bugs in complex codes”. In: *Proceedings of the International Conference on High Performance Computing, Networking, Storage and Analysis*. ACM. 2013, p. 45.
- [150] Aniello Castiglione et al. “Exploiting mean field analysis to model performances of big data architectures”. In: *Future Generation Computer Systems* 37 (2014). Special Section: Innovative Methods and Algorithms for Advanced Data-Intensive Computing Special Section: Semantics, Intelligent processing and services for big data Special Section: Advances in Data-Intensive Modelling and Simulation Special Section: Hybrid Intelligence for Growing Internet and its Applications, pp. 203–211. ISSN: 0167-739X. DOI: <https://doi.org/10.1016/j.future.2013.07.016>. URL: <http://www.sciencedirect.com/science/article/pii/S0167739X13001611>.
- [151] HT Intema et al. “The GMRT 150 MHz all-sky radio survey-First alternative data release TGSS ADR1”. In: *Astronomy & Astrophysics* 598 (2017), A78.
- [152] S Kazemi et al. “Radio interferometric calibration using the SAGE algorithm”. In: *Monthly Notices of the Royal Astronomical Society* 414.2 (2011), pp. 1656–1666.
- [153] Ronny Levanda and Amir Leshem. “Synthetic aperture radio telescopes”. In: *Signal Processing Magazine, IEEE* 27 (Feb. 2010), pp. 14–29. DOI: 10.1109/MSP.2009.934719.
- [154] Cecchi Marco et al. “The glite workload management system”. In: *Journal of Physics: Conference Series*. Vol. 219. IOP Publishing. 2010, p. 062039.
- [155] I Virtanen et al. *Station Data Cookbook v1.2*. 2018. URL: [http://lofar.ie/wp-content/uploads/2018/03/station\\_data\\_cookbook\\_v1.2.pdf](http://lofar.ie/wp-content/uploads/2018/03/station_data_cookbook_v1.2.pdf).
- [156] Eric Jones, Travis Oliphant, Pearu Peterson, et al. *SciPy: Open source scientific tools for Python*. [Online; accessed November 21, 2019]. 2001–. URL: <http://www.scipy.org/>.
- [157] Huub Rottgering et al. “LOFAR and APERTIF Surveys of the Radio Sky: Probing Shocks and Magnetic Fields in Galaxy Clusters”. English. In: *Journal of Astrophysics and Astronomy* 32.4 (Dec. 2011), pp. 557–566. ISSN: 0250-6335.

- [158] Gregory M. Kurtzer, Vanessa Sochat, and Michael W. Bauer. “Singularity: Scientific containers for mobility of compute”. In: *PLOS ONE* 12.5 (May 2017), pp. 1–20. DOI: 10.1371/journal.pone.0177459. URL: <https://doi.org/10.1371/journal.pone.0177459>.
- [159] Martin Fowler and Matthew Foemmel. “Continuous integration”. In: *Thought-Works* <http://www.thoughtworks.com/Continuous Integration.pdf> 122 (2006), p. 14.
- [160] Jez Humble and David Farley. *Continuous Delivery: Reliable Software Releases through Build, Test, and Deployment Automation (Adobe Reader)*. Pearson Education, 2010.
- [161] Andrew E. Slaughter et al. “Continuous integration for concurrent MOOSE framework and application development on GitHub”. In: *Journal of Open Research Software* 3.1 (Nov. 2015). ISSN: 2049-9647. DOI: 10.5334/jors.bx.
- [162] Pasquale Salza, Filomena Ferrucci, and Federica Sarro. “Develop, deploy and execute parallel genetic algorithms in the cloud”. In: *Proceedings of the 2016 on Genetic and Evolutionary Computation Conference Companion*. ACM, 2016, pp. 121–122.
- [163] Matthew C Chambers et al. “A cross-platform toolkit for mass spectrometry and proteomics”. In: *Nature biotechnology* 30.10 (2012), p. 918.
- [164] Eric Jeschke and Takeshi Inagaki. “Lessons learned deploying a second generation Observation Control System for Subaru Telescope”. In: *Software and Cyberinfrastructure for Astronomy*. Vol. 7740. International Society for Optics and Photonics, 2010, 77400F.
- [165] *Travis CI - Test and Deploy Your Code with Confidence*. <https://travis-ci.org/>. Accessed: 2019-04-18.
- [166] *Continuous Integration and Delivery - CircleCI*. <https://circleci.com/>. Accessed: 2019-04-18.
- [167] *GitLab Continuous Integration & Delivery | GitLab*. <https://about.gitlab.com/product/continuous-integration/>. Accessed: 2019-04-18.
- [168] *Jenkins*. <https://jenkins.io/>. Accessed: 2019-04-18.
- [169] Davide Salomoni et al. “INDIGO-Datacloud: foundations and architectural description of a Platform as a Service oriented to scientific computing”. In: *CoRR* abs/1603.09536 (2016). arXiv: 1603.09536. URL: <http://arxiv.org/abs/1603.09536>.
- [170] Carl Boettiger. “An Introduction to Docker for Reproducible Research”. In: *SIGOPS Oper. Syst. Rev.* 49.1 (Jan. 2015), pp. 71–79. ISSN: 0163-5980. DOI: 10.1145/2723872.2723882. URL: <http://doi.acm.org/10.1145/2723872.2723882>.

- [171] Oscar Esteban et al. “fMRIPrep: a robust preprocessing pipeline for functional MRI”. In: *Nature methods* 16.1 (2019), p. 111.
- [172] Björn Grüning et al. “Practical computational reproducibility in the life sciences”. In: *Cell systems* 6.6 (2018), pp. 631–635.
- [173] Andrew J Younge et al. “A tale of two systems: Using containers to deploy hpc applications on supercomputers and clouds”. In: *2017 IEEE International Conference on Cloud Computing Technology and Science (CloudCom)*. IEEE. 2017, pp. 74–81.
- [174] Yaron Goland et al. *HTTP Extensions for Distributed Authoring–WEBDAV*. Tech. rep. 1999.
- [175] Stefan Van Der Walt, S. Chris Colbert, and Gaël Varoquaux. “The NumPy array: a structure for efficient numerical computation”. In: *Computing in Science & Engineering* 13, arXiv:1102.1523 (2 Mar. 2011), pp. 22–30. DOI: 10.1109/MCSE.2011.37. arXiv: 1102.1523 [cs.MS].
- [176] The Astropy Collaboration et al. “The Astropy Project: Building an inclusive, open-science project and status of the v2.0 core package”. In: *ArXiv e-prints* (Jan. 2018). arXiv: 1801.02634 [astro-ph.IM].
- [177] Gary J Hill et al. “The Hobby-Eberly Telescope Dark Energy Experiment (HETDEX): Description and Early Pilot Survey Results”. In: *arXiv preprint arXiv:0806.0183* (2008).
- [178] Mark D Wilkinson et al. “The FAIR Guiding Principles for scientific data management and stewardship”. In: *Scientific data* 3 (2016).
- [179] Melanie Johnston-Hollitt. “Taming the Data Deluge to Unravel the Mysteries of the Universe”. In: *Proceedings of the 26th International Conference on World Wide Web*. International World Wide Web Conferences Steering Committee. 2017, pp. 1–1.
- [180] **Alexandar P Mechev** et al. “Automated testing and quality control of LOFAR scientific pipelines with AGLOW”. In: *Astronomy and Computing* (2019 in review).
- [181] T. W. Shimwell et al. “The LOFAR Two-metre Sky Survey - II. First data release”. In: *A&A* 622 (2019), A1. DOI: 10.1051/0004-6361/201833559. URL: <https://doi.org/10.1051/0004-6361/201833559>.
- [182] K. L. Emig et al. “Searching for the largest bound atoms in space”. In: *A&A* (2019).
- [183] Y. G. Grange et al. “Characterising radio telescope software with the Workload Characterisation Framework”. In: *arXiv e-prints*, arXiv:1612.00456 (Dec. 2016), arXiv:1612.00456. arXiv: 1612.00456 [astro-ph.IM].
- [184] **A.P. Mechev** et al. “Building LOFAR as a Service”. In: *Astronomical Data Analysis Software and Systems XXVIII*, p. 677.