

Can Clouds Replace Grids? Will Clouds Replace Grids?

J.D. Shiers

CERN, 1211 Geneva 23, Switzerland

Jamie.Shiers@cern.ch

Abstract. The world's largest scientific machine – comprising dual 27km circular proton accelerators cooled to 1.9°K and located some 100m underground – currently relies on major production Grid infrastructures for the offline computing needs of the 4 main experiments that will take data at this facility. After many years of sometimes difficult preparation the computing service has been declared “open” and ready to meet the challenges that will come shortly when the machine restarts in 2009. But the service is not without its problems: reliability – as seen by the experiments, as opposed to that measured by the official tools – still needs to be significantly improved. Prolonged downtimes or degradations of major services or even complete sites are still too common and the operational and coordination effort to keep the overall service running is probably not sustainable at this level. Recently “Cloud Computing” – in terms of pay-per-use fabric provisioning – has emerged as a potentially viable alternative but with rather different strengths and no doubt weaknesses too. Based on the concrete needs of the LHC experiments – where the total data volume that will be acquired over the full lifetime of the project, including the additional data copies that are required by the Computing Models of the experiments, approaches 1 Exabyte – we analyze the pros and cons of Grids versus Clouds. This analysis covers not only technical issues – such as those related to demanding database and data management needs – but also sociological aspects, which cannot be ignored, neither in terms of funding nor in the wider context of the essential but often overlooked role of science in society, education and economy.

1. Introduction

The first In order to process and analyze the data from the world's largest scientific machine, a worldwide grid service – the Worldwide LHC Computing Grid (LCG) [1] – has been established, building on two main production infrastructures: those of the Open Science Grid (OSG) [2] in the Americas, and the Enabling Grids for E-sciencE (EGEE) [3] Grid in Europe and elsewhere. The machine itself – the Large Hadron Collider (LHC) – is situated some 100m underground beneath the French-Swiss border near Geneva, Switzerland and supports four major collaborations and their associated detectors: ATLAS, CMS, ALICE and LHCb.

Even after several levels of reduction, some 15PB of data will be produced per year at rates to persistent storage of up to 1.5GB/s – the LHC itself having an expected operating lifetime of some 10 – 15 years. These data will be analyzed by scientists at close to two hundred and fifty institutes worldwide using the distributed services that form the Worldwide LHC Computing Grid (WLCG)[4][5][6]. Depending on the computing models of the various experiments, additional data copies are made at the various institutes, giving a total data sample well in excess of 500PB and possibly exceeding 1EB.

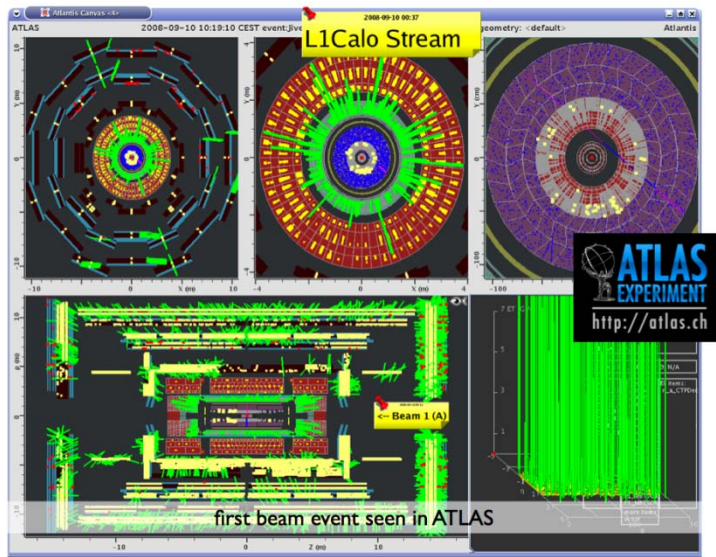


Figure 1 - First Beam Event Seen in the ATLAS Detector

Running a service where the user expectation is for support 24x7, with extremely rapid problem determination and resolution targets, is already a challenge. When this is extended to a large number of rather loosely coupled sites, the majority of which support multiple disciplines – often with conflicting requirements but always with local constraints – this becomes a major or even “grand” challenge. That this model works at the scale required by the LHC experiments – literally around the world and around the clock – is a valuable vindication of the Grid computing paradigm.

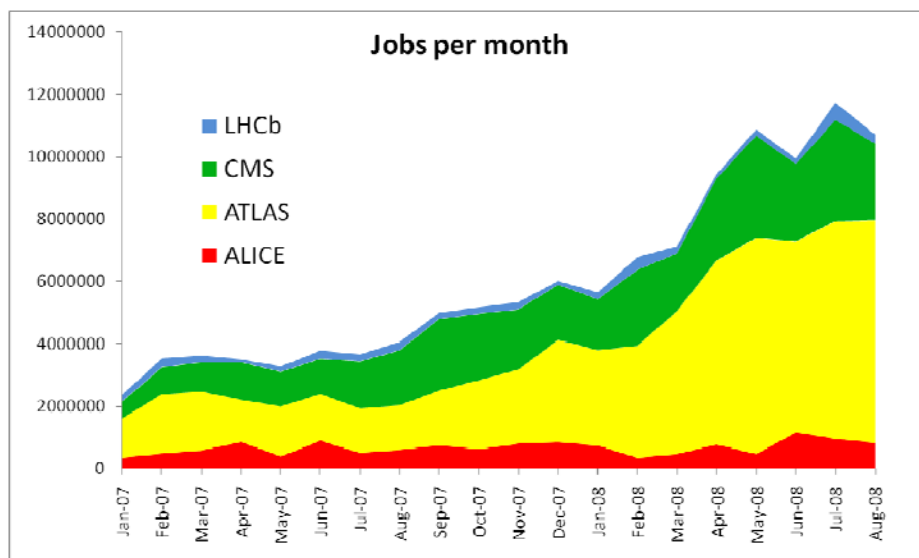


Figure 2 - Jobs per month by LHC virtual organisation

However, even after many years of preparation – including the use of well-proven techniques for the design, implementation, deployment and operation of reliable services – the operational costs are still too high to be sustained in the long term. This translates to significant user frustration and even disillusionment. On the positive side, however, the amount of application support that is required compares well with that of some alternate models, such as those based on supercomputers. The costs involved with such solutions are way beyond the means of the funding agencies involved, nor are they

necessarily well adapted to the “embarrassingly parallel” nature of the types of data processing and analysis that typify the High Energy Physics (HEP) domain.

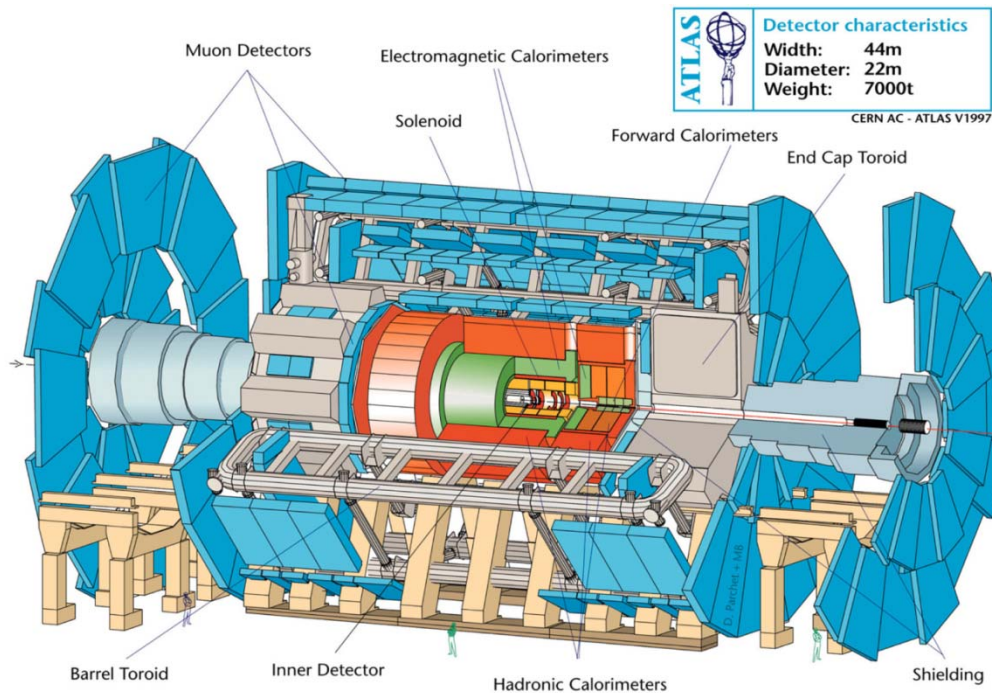


Figure 3 - The ATLAS Detector

This makes HEP an obvious test-case for cloud computing models and indeed a number of feasibility studies have already been performed. The purpose of this paper is to explore the potential use of clouds against a highly ambitious target: not simply whether it is possible on paper – or even in practice – to run applications that typify our environment, but whether it would be possible and affordable to deliver a level of service equivalent to – or even higher – than that available today using Grid solutions. In addition to analyzing the technical challenges involved, the “hidden benefits” of Grid computing, namely in terms of the positive feedback provided – both scientifically and culturally – to the local institutes and communities that provide resources to the Grid, and hence to their funding agencies who are thus hopefully motivated to continue or even increase their level of investment, are also compared. Finally, based on the wide experience gained by sharing Grid solutions with a large range of disciplines, we try to generalize these findings to make some statements regarding the benefits and weaknesses of these competing – or possibly simply complementary – models.

2. Motivation

There is a wide range of applications that require significant computational and storage resources – often beyond what can conveniently be provided at a single site. These applications can be broadly categorized as *provisioned* – meaning that the resources are needed more or less continuously for a period similar to, or exceeding, the usable lifetime of the necessary hardware; *scheduled* – where the resources are required for shorter periods of time and the results are not necessarily time critical (but higher than for the following category); *opportunistic* – where there is no urgent time pressure, but any available resources can be readily soaked up. Reasons why the resources cannot easily be provided at a single site include those of funding, where international communities are under pressure to spend funds locally to institutes that are part of the collaboration, as well as those of power and cooling – increasingly a problem with high energy prices and concerns over greenhouse gases.

Whilst Grid computing can claim significant successes in handling the needs of these communities and their applications, the entry threshold – both for new applications as well as additional sites / service providers – is still considered too high and is an impediment to their wide-scale adoption. Nevertheless, one cannot deny the importance of many of the applications currently investigating or using Grid technologies, including drug research, disaster response and prediction, as well as major scientific research areas, typified by High Energy Physics and CERN’s Large Hadron Collider programme, amongst many others.

Currently, adapting an existing application to the Grid environment is a non-trivial exercise that requires an in-depth understanding not only of the Grid computing paradigm but also of the computing model of the application in question.

The successful demonstration of a straightforward recipe for moving a wide range of applications – from simple to the most demanding – to Cloud environments would be a significant boost for this technology and could open the door to truly ubiquitous computing. This would be similar to the stage when the Web burst out of the research arena and use by a few initiates to its current state as a tool used by virtually everyone as part of their everyday work and leisure. However, the benefits can be expected to be much greater – given that there is essentially unlimited freedom in the type of algorithms and volumes of data that can be processed.

3. Service Targets

There are two distinct views of the service targets for WLCG: those specified up-front in a Memorandum of Understanding [7] (MoU) – signed by the funding agencies that provide the resources to the Grid – and the “expectations” from the experiments. We have seen a significant mismatch between these two views and attempted to reconcile them into a single set of achievable and measurable targets.

Service	Maximum delay in responding to problems		
	Interruption	Degradation > 50%	Degradation > 20%
Raw data recording	4 hours	6 hours	6 hours
Event reconstruction or distribution of data to Tier1s	6 hours	6 hours	12 hours
Networking service to Tier1s	6 hours	6 hours	12 hours

Table 1- Extract of Service Targets (Tier0)

The basic underlying principle is not to “guarantee” perfect services, but to focus on specific failure modes, limit them where possible, and ensure sufficient redundancy is built in at the required levels to allow “automatic” recovery from failures – e.g. using buffers and queues of sufficient size that are automatically drained once the corresponding service is re-established. Nevertheless, the targets remain high, specified both in service availability measured on an annual basis as well as the time to respond when necessary.

Criticality of Service	Impact of degradation / loss
Very high	interruption of these services <u>affects online data-taking operations or stops any offline operations</u>
High	interruption of these services <u>perturbs seriously offline computing operations</u>
Moderate	interruption of these services perturbs software development and part of computing operations

Table 2- Service Criticality (ATLAS virtual organisation)

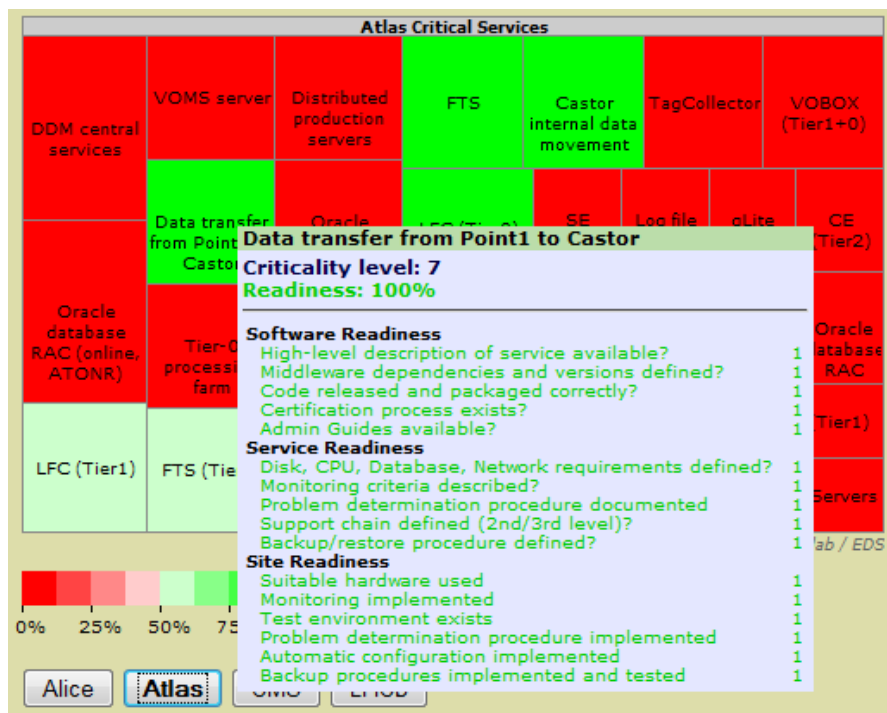


Figure 4- "GridMap" visualization of service readiness

Software readiness
High-level description of service available?
Middleware dependencies and versions defined?
Code released and packaged correctly?
Certification process exists?
Admin Guides available?
Service readiness
Disk, CPU, Database, Network requirements defined?
Monitoring criteria described?
Problem determination procedure documented?
Support chain defined (2nd/3rd level)?
Backup/restore procedure defined?
Site readiness
Suitable hardware used?
Monitoring implemented?
Test environment exists?
Problem determination procedure implemented?
Automatic configuration implemented?
Backup procedures implemented and tested?

Table 3 - Service Readiness Checklist

As currently defined, a very small number of incidents are sufficient to bring a site below its availability target. In order to bridge the gap between these two potentially conflicting views – and building on the above mentioned industry-standard techniques – we observe relatively infrequent breaks of service: either those that are directly user-visible or those that cannot be smoothed over using the buffering and other mechanisms mentioned. We have put in place mechanisms whereby appropriately privileged members of the user communities can raise alarms in these case –

supplementing the automatic monitoring that may not pick up all error conditions – that can be used 24x7 to alert the support teams at a given site.

Whilst these mechanisms are used relatively infrequently – around once per month in the most intense periods of activity – the number of situations where a major service or site is either degraded or unavailable for prolonged periods of time with respect to the targets defined in the MoU is still far too high – sometimes several times per week. Most of these failures fall into a small number of categories:

- Power and cooling – failures in a site’s infrastructure typically have major consequences – the site is down for many hours. Whilst complete protection against such problems is unlikely to be affordable, definition and testing of recovery procedures could be improved – e.g. ensuring the order in which services are restarted is well understood and adhered to, making sure that the necessary infrastructure – redundant power supplies, network connections and so forth – are such as to maximize protection and minimize the duration of any outages;
- Configuration issues – required configuration changes are often communicated in a variety of (unsuitable) formats, with numerous transcription (and even interpretation) steps, all sources of potential errors;
- Database and data management services – the real killers. For our data intensive applications, these typically render a site or even region unusable.

Rank	Services at Tier0
Very high	Oracle (online), DDM central catalogues
High	P1→T0 transfers, online-offline DB connectivity, CASTOR internal data movement, T0 processing farm, Oracle (offline), LFC, FTS, VOMS, Dashboard, Panda/Bamboo, DDM site services
Moderate	3D streaming, WMS, SRM/SE, CAF, CVS, AFS, build system
Rank	Services at Tier1
High	LFC, FTS
Moderate	3D streaming, Oracle, SRM/SE, CE
Rank	Services at Tier2
	SRM/SE, CE

Table 4 - Ranking of ATLAS services across main tiers

The above table (a detailed description of all of the acronyms is not relevant here but can be found in [5]) emphasizes the importance of database and data management services: the most and second most critical services required at the Tier0 and Tier1 sites are either database related, data management related, or in most cases both.

When	Issue	Target
Now	Consistent use of all Service Standards	100%
30'	Operator response to alarm / alarm e-mail	99%
1 hour	Operator response to alarm / alarm e-mail	100%
4 hours	Expert intervention in response to above	95%
8 hours	Problem resolved	90%
24 hours	Problem resolved	99%

Table 5 - Targets for Tier0 Services

These service targets are complemented by more specific requirements from the experiments. The tables below list those for the CMS experiment. A site can be in one of the following 3 states:

1. **COMMISSIONED:** *daily rules satisfied during the last 2 days, or during the last day and at least 5 days in the last 7*

2. **WARNING:** daily rules not satisfied in the last day but satisfied during at least 5 days in the last 7
3. **UNCOMMISSIONED:** daily rules satisfied for less than 5 days in the last 7

The purpose of these rules is to ensure as many sites as possible stay in commissioned status and to allow for a fast recovery when problems start to occur.

Daily Rules for Tier1 sites
Daily SAM (service availability monitoring) $\geq 90\%$
Daily job robot efficiency) $\geq 95\%$
Having commissioned the downlink with the Tier0
Having ≥ 10 commissioned downlinks to the Tier2 sites
Having ≥ 4 commissioned down/uplinks to other Tier1 sites

Table 6 – CMS targets for Tier1 sites

Daily Rules for Tier2 sites
Daily SAM (service availability monitoring) $\geq 80\%$
Daily job robot efficiency) $\geq 90\%$
Having a commissioned uplink with at least 1 Tier1
Having a commissioned downlink with ≥ 2 Tier1 sites

Table 7 - CMS targets for Tier-2 sites

The following figure shows a historical snap-shot of CMS Tier2 sites for the specified time-window.

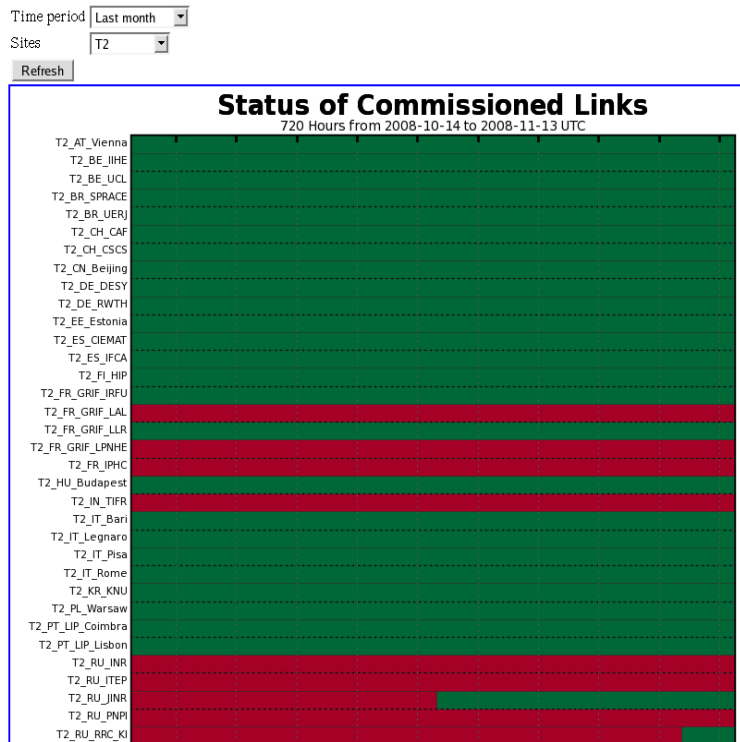


Figure 5 - Status of CMS Links

In principle, grids – like clouds – should offer sufficient redundancy that the failure of some fraction of the overall system can be tolerated with little or preferably no service impact. This is not, unfortunately, true of all computing models in use in HEP, in which for reasons of both geography and funding specific dependencies exist between different sites – both nationally and internationally.

Furthermore, sites have well defined functional roles in the overall data processing and analysis chain which mean that they cannot always be replaced by any other – although sometimes by one or more specific sites. This is not a weakness of the underlying model but simply a further requirement from the application domain – the proposed solution must also work given the requirements and constraints from the possibly sub-optimal computing model involved.

The main responsibilities of the first 3 tiers are given below:

- Tier0 (CERN): safe keeping of RAW data (first copy); first pass reconstruction, **distribution of RAW data and reconstruction output (Event Summary Data or ESD) to Tier1**; reprocessing of data during LHC down-times;
- Tier1: safe keeping of a proportional share of RAW and reconstructed data; large scale **reprocessing** and safe keeping of corresponding output; **distribution of data products to Tier2s** and safe keeping of a share of simulated data produced at these Tier2s;
- Tier2: Handling **analysis** requirements and proportional share of **simulated event** production and reconstruction.

In the considerations below, we will discuss not only whether the cloud paradigm could be used to solve all aspects of LHC computing but also whether it could be used for the roles provided by one or more tiers or for specific functional blocks (e.g. analysis, simulation, re-processing etc.)

4. The data is the challenge

Whilst there is little doubt that for applications that involve relatively small amounts of data and/or data rates the cloud computing model is almost immediately technically viable this is one of the largest areas of concern for our application domain. Specific issues include:

- Long-term data curation: if this is the responsibility of “the user” a significant amount of infrastructure and associated support is required to store and periodically migrate data between old and new technologies over long periods of time – problems familiar to those involved with large scale (much more than 1PB) data archives;
- Data placement and access: although we have been relatively successful in defining standard interfaces to a reasonably wide-range of storage system implementations rather fine-grained control on data placement and data access has been necessary to obtain the necessary performance and isolation of the various activities – both between and within virtual organizations;
- Data transfer: possibly a curiosity of the computing models involved and strongly coupled to the specific roles of the sites that make up the WLCG infrastructure – bulk data currently needs to be transferred at high rates in pseudo real-time between sites. Would this be simplified or eliminated using a cloud-based solution? The figure below shows the percentage of file transfers that are successful on the first attempt. It is clearly much lower than desirable, resulting in wasted network bandwidth and extra load on the storage services, which in turn has a negative effect on other activities;
- Database applications: behind essentially all data management applications even if a variety of technologies are used – often at a single site. Again, deep knowledge of the hardware configuration and physical implementation are currently required to get an acceptable level of service.

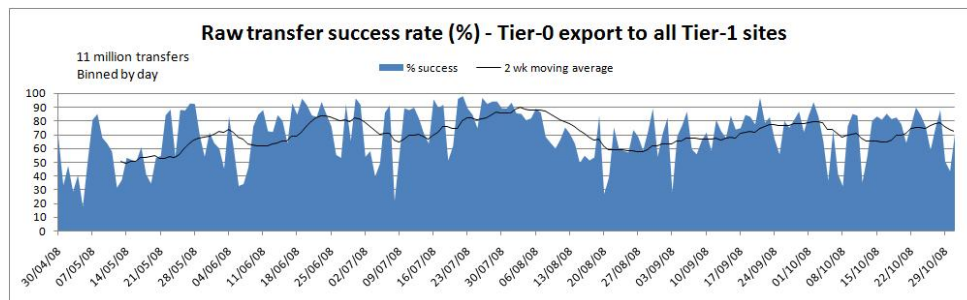


Figure 6 - Success Rate of File Transfers

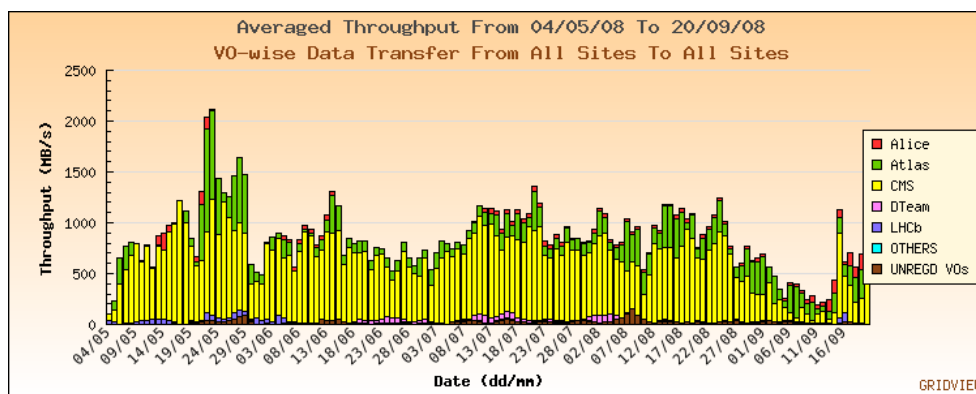


Figure 7 – Averaged Daily Data Transfer Rates

It is perhaps unfair to compare a solution that has evolved over around a decade, with many teething problems and a number of major outstanding issues, with an alternative and impose that a well-established computing model must be supported without change. On the other hand, targets that are relatively independent of the implementation can be defined in terms of availability, service level, computational and data requirements. It may well be that on balance the technical and managerial advantages outweigh any as yet to be found drawbacks. This would leave unavoidable issues such as cost, together with the sociological and other “spin-off” benefits.

5. Operations Costs

The operations costs of a large-scale grid infrastructure are rarely reported and even when this is done it typically refers to the generic infrastructure and not to the total costs of operating the computing infrastructure of a large-scale collaboration. At neither level are the costs negligible: the European Grid Initiative Design Study estimates the total number of full time equivalents for operations-related activities across all National Grid Initiatives to be broadly in the 200-400 range – with an extremely modest 5(!) people providing overall coordination (compared to 15-20 in the EGEE project, currently in its 3rd and presumably final phase). The operations effort required for a single large virtual organization, such as the ATLAS experiment – the largest LHC collaboration, is almost certainly in excess of 100. A typical WLCG “Grid Deployment Board” – the monthly meeting working on the corresponding issues – also involves around 100 local and remote participants, whereas a “WLCG Collaboration workshop” can attract closer to 300 – mainly site administrators and other support staff. These costs are not always easy to report accurately as they often covered – at least in part – by doctoral students, post-doctoral fellows and other “dark effort”. However, any objective comparison between different solutions must include the total cost of ownership and not just a somewhat arbitrary subset.

6. Grid-Based Petascale Production is Reality

Despite the remaining rough edges to the service, as well as the undeniably high operational costs, the success of building a petascale (using the loose definition of 100,000 cores) world-wide distributed production facility that is built using several independently managed and funded major grid infrastructures – of which the two main components are built out of $O(100)$ sites (EGEE and OSG) – must be considered a large success. The system has been in production mode – with steady improvement in reliability over time – since at least 2005. This includes formal capacity planning, scheduling of interventions – the majority of which can be performed with zero user-visible downtime – and regular reviews of availability and performance metrics. A service capable of meeting the evolving requirements of the LHC experiments must continue for at least the usable life of the accelerator itself plus an additional few years for the main analysis of the data to be completed. Including foreseen accelerator and detector upgrades, this probably means until around 2030! Whether grids survive this long is somewhat academic – major changes in IT are inevitable on this timescale and adapting to (or rather benefiting from) these advances are required. How this can be done in a non-disruptive manner is certainly a challenge but it is worth recalling that the experiments at LEP – the previous collider in the same tunnel – started in an almost purely mainframe environment (IBM, Cray, large VAXclusters and some Apollo workstations) and moved first to farms of powerful Unix workstations (HP, SGI, Sun, IBM, ...) and finally PCs running Linux. This was done without interruption to on-going data taking, reprocessing and analysis, but obviously not without major work. Much more recently, several hundred TB of data from several experiments went through a “triple migration” – a change of backend tape media, a new persistency solution and a corresponding re-write of the offline software – a major effort involving many months of design and testing and an equivalent period for the data migration itself. (The total effort was estimated at $\sim 1\text{FTE}/100\text{TB}$ of data migrated). These examples give us confidence that we are able to adapt to major changes of technology that are simply inevitable for projects with lifetimes measured in decades.

7. Towards a Cloud Computing Challenge

Over the past few years a series of “service challenges” has been carried out to ramp-up to the necessary level required to support data taking and processing at the LHC. This culminated in 2008 in a so-called “Common Computing Readiness Challenge (CCRC’08)” – aimed at showing that the computing infrastructure was ready to meet the needs of all supported experiments at all sites that support them. Given a large number of changes foreseen prior to data taking in 2009, a further “CCRC’09” is scheduled for 2 months prior to data taking in that year. (This will be a rather different event than the 2008 challenge, relying on on-going production activities, rather than scheduled tests, to generate the necessary workload. Where possible, overlap of inter-VO, as well as intra-VO, activities will be arranged to show that the system can handle the combined workloads satisfactorily). An important – indeed necessary – feature of these challenges has been metrics that are agreed upfront and are reported on regularly to assess our overall state and progress. Whilst it is unlikely that in the immediate future a challenge on an equivalent scale could be performed using a cloud environment, such a demonstration is called for – possibly at progressively increasing scale – if the community is to be convinced of the validity and even advantages of such an approach.

The obvious area where to start is that of simulation – a compute-dominated process with relatively little I/O needs. Furthermore, in the existing computing models the Tier2 sites – where such work typically but not exclusively takes place – data curation is not provided. Thus, the practice of storing output data at a (Tier1) site that does provide such services is well established. Thus, the primary question that should be answered by such a question is:

- Can cloud computing offer compute resources for low I/O applications, including services for retrieval of output data for long-term data storage “outside” of the cloud environment, in a manner that is sufficiently performant as well as cost-competitive with those typically offered today by Universities and smaller institutes?

To perform such a study access to the equivalent of several hundred – a few thousand cores for a minimum of some weeks would be required. There is little doubt that such a study would be successful from a technical point of view, but would it be not only competitive or even cheaper in terms of total cost of ownership? The requirements for such a study have been oversimplified – e.g. the need for access to book-keeping systems and other database applications and a secure authentication mechanism for the output storage – but it would make a valuable first step. If not at least in the same ball-park in terms of the agreed criteria there would be little motivation for further studies.

Rather than loop through the various functional blocks that are mapped to the various tiers described above, further tests could be defined in terms of database and data management functionality – presumably both more generic as well as more immediately understandable to other disciplines. These could be characterized in terms of the number of concurrent streams, the type and frequency of access (sequential, random, rarely, frequent) and equivalent criteria for database applications. These are unlikely to be trivial exercises but the potential benefit is large – one example being the ability of a cloud-base service to adapt to significant changes in needs, such as pre-conference surges that can typically not be accommodated by provisioned resources that do not have enough headroom for such peaks, often synchronized across multiple activities, both within and across multiple virtual organizations.

8. Data Grids and Computational Clouds – Friends or Foes?

The possibility of Grid computing taking off in a manner somehow analogous to that of the Web has often been debated. A potential stumbling block has always been cost and subscription models analogous to those of mobile phone network providers have been suggested. In reality, access to the Web is often not “free” – there may not be an explicit charge for Internet access in many companies and institutes – and without the Internet the Web would have little useful meaning. However, for most people Internet access is through a subscription service, that may itself be bundled with others, such as “free” national or even international phone calls, access to numerous TV channels and other such services.

A more concrete differentiator is the “closed” environment currently offered as “Clouds” – it may be clear how one purchases services but not how one contributes computational and storage resources in the manner that a site can “join” an existing Grid. A purely computational Grid – loosely quantified as one that provides no long term data storage facilities or curation – is perhaps the most obvious competitor of Clouds. Assuming such facilities are shared as described above between provisioned, scheduled and opportunistic use a more important distinction could – again – be in the level of data management and database services that are provided.

A fundamental principle of our grid deployment model has been to specify the interfaces but not the implementation. This has allowed sites to accommodate local requirements and constraints whilst still providing interoperable services. It has, however, resulted in a much higher degree of complexity and in less pooling of experience and techniques than could otherwise have been the case. This is illustrated when the strategies for two of the key components – databases and data management – are compared. The main database services at the Tier0 and Tier1 sites (at least for ATLAS – the largest VO), have been established using a single technology (Oracle) with common deployment and operational models. Data management services – whilst accessed through a common interface (SRM) – are implemented in numerous different variations. Even when the same software solution is used, the deployment model differs widely and it has proven hard to share experience. The table below shows the diversity in terms of front-end storage solutions: in the case of dCache not only are multiple releases deployed but also the backend tape-based mass storage system (both hardware and software) varies from site to site – creating additional complexity.

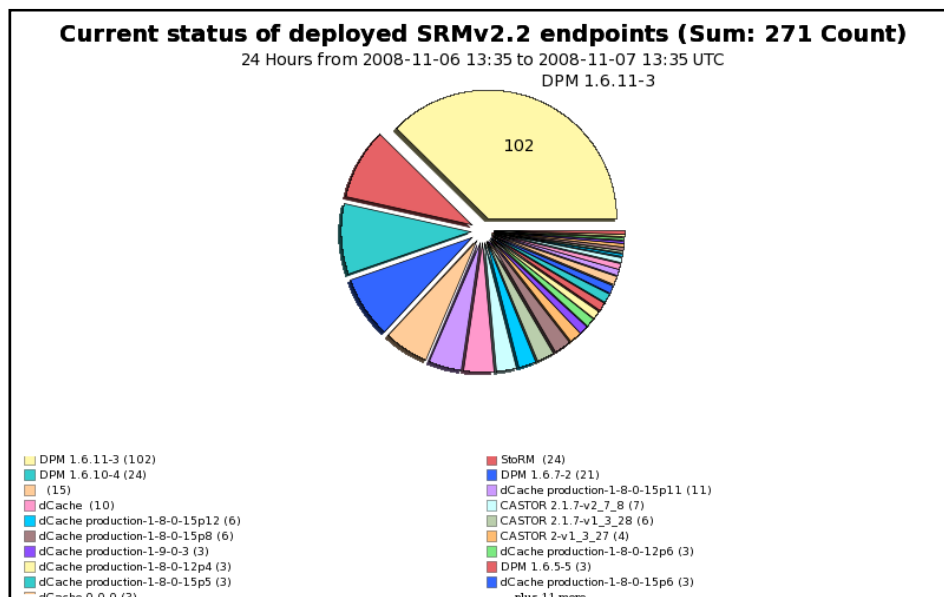


Figure 8 - Distribution of Storage Solutions and Versions

There is little doubt that the cost of providing such services as well as the achieved service level suffers as a result – even if more “politically correct”. Any evolution or successor of these services would benefit from learning from these experiences.

9. Grids versus Clouds – Sociological Factors

For many years an oft-levelled criticism of HEP has been the “brain-drain” effect from Universities and other institutes to large central facilities such as CERN. Although distributed computing has been in place since before the previous generation of experiments at the LEP collider – formerly housed in the same 27km tunnel as the LHC today – scientists at the host laboratory had very different possibilities to those at regional centres or local institutes. Not only does the grid devolve extremely important activities to the Tier1 and Tier2 sites but the key question of equal access to all of the data is essentially solved. This brings with it the positive feedback effect mentioned above which is so important that it probably outweighs even a (small) cost advantage – to be proven – in favour of non-grid models.

10. Is the Gain worth the Pain?

It should be clear from the above that some of the major service problems associated with today’s production grid environment could be avoided by adopting a simpler deployment model: fewer sites, less diversity but also less flexibility. However, much of the funding that we depend on would not be readily available unless it was spent – as now – primarily locally. On the other hand – and in the absence of any large-scale data-intensive tests – it is unclear whether a cloud solution could meet today’s technical requirements. A middle route is perhaps required, whereby grid service providers learn from the difficulties and costs of providing reliable but often heterogeneous services, as well as the advantages in terms of service level, possibly at the cost of some flexibility, through a more homogeneous approach. Alternatively, some of the peak load could perhaps be more efficiently and cost effectively handled by cloud computing, leaving strongly data-related issues to the communities that own them and are therefore presumably highly motivated to solve them.

For CERN, answers to these questions are highly relevant – projections show that we will run out of power and cooling in the existing computer centre on a time-scale that precludes building a new one on the CERN site (for obvious reasons, priority has been given in recent years to the completion of the

LHC machine). Overflow capacity maybe available in a partner site to tide us through: do we have the time to perform a sufficiently large scale demonstration of a cloud-based solution to obviate such a move? Is there a provider sufficiently confident of their solution that they are willing to step up to this challenge? There have been no takers so far and time is running out – at least for this real-life exabyte-scale test-case. In the meantime our focus is on greatly improving the stability and usability of our storage services, not only to handle on-going production activity with acceptably low operational costs, whilst preparing for large-scale data-intensive end-user analysis that will come with the first real data from the world's largest scientific machine.

11. Conclusions

After many years of research and development followed by production deployment and usage by many VOs, worldwide Grids that satisfy the criteria in Ian Foster's "grid checklist" [8] are a reality. There is significant interest in longer-term sustainable infrastructures that are compatible with the current funding models and work on the definition of the functions of and funding for such systems is now underway. Using a very simple classification of Grid applications, we have briefly explored how the corresponding communities could share common infrastructures to their mutual benefit. A major challenge for the immediate future is the containment of the operational and support costs of Grids, as well as reducing the difficulties in supporting new communities and their applications. These and other issues are being considered by a design study for a long term e-infrastructure [9]. Cloud computing may well be the next step in the long road from extremely limited computing – as typified by the infamous Thomas J. Watson 1943 quote "*I think there is a world market for maybe five computers*" – to a world of truly ubiquitous computing (which does not mean free). It is clear that the applications described in this document may represent today's "lunatic fringe", but history has repeatedly shown that these needs typically become main-stream within only a few years. We have outlined a number of large-scale production tests that would need to be performed in order to assess clouds as complementary or even replacement technology for the grid-based solutions in use today, although data-related issues remain a concern. Finally, we have raised a number of non-technical, non-financial concerns that must nevertheless be taken into account – particularly by large-scale research communities that rely on various funding sources and must – for their continued existence – show value to those that ultimately support them: sometimes a private individual or organization but often the tax-payer.

References

- [1] The Worldwide LHC Computing Grid (WLCG), <http://lcg.web.cern.ch/LCG/>.
- [2] The Open Science Grid, <http://www.opensciencegrid.org/>.
- [3] The Enabling Grids for E-science (EGEE) project, <http://public.egee.org/>.
- [4] J. D. Shiers, "*Lessons Learnt from Production WLCG Deployment*", to appear in the proceedings of the International Conference on Computing in High Energy Physics, Victoria, BC, September 2007.
- [5] The Worldwide LHC Computing Grid (worldwide LCG), Computer Physics Communications 177 (2007) 219–223, Jamie Shiers, CERN, 1211 Geneva 23, Switzerland
- [6] LCG Technical Design Report, CERN-LHCC-2005-024, available at <http://lcg.web.cern.ch/LCG/tdr/>.
- [7] Memorandum of Understanding for Collaboration in the Deployment and Exploitation of the Worldwide LHC Computing Grid, available at <http://lcg.web.cern.ch/LCG/C-RRB/MoU/WLCGMoU.pdf>.
- [8] I. Foster, Argonne National Laboratory and University of Chicago, What is the Grid? A Three Point Checklist, 2002.
- [9] The European Grid Initiative Design Study – website at <http://www.eu-egi.org/>.