

e-Science and its implications

BY TONY HEY[†] AND ANNE TREFETHEN

*UK e-Science Core Programme,
Engineering and Physical Sciences Research Council,
Polaris House, Swindon SN2 1ET, UK*

Published online 25 June 2003

After a definition of e-science and the Grid, the paper begins with an overview of the technological context of Grid developments. NASA's Information Power Grid is described as an early example of a 'prototype production Grid'. The discussion of e-science and the Grid is then set in the context of the UK e-Science Programme and is illustrated with reference to some UK e-science projects in science, engineering and medicine. The Open Standards approach to Grid middleware adopted by the community in the Global Grid Forum is described and compared with community-based standardization processes used for the Internet, MPI, Linux and the Web. Some implications of the imminent data deluge that will arise from the new generation of e-science experiments in terms of archiving and curation are then considered. The paper concludes with remarks about social and technological issues posed by Grid-enabled 'collaboratories' in both scientific and commercial contexts.

Keywords: e-Science; Grid; data curation; data preservation; open standards

1. Introduction

The term 'e-science' was introduced by John Taylor, Director General of Research Councils in the UK Office of Science and Technology (OST). From his previous experience as Head of Hewlett-Packard's European Research Laboratories in Bristol, and from his experience in his current post, Taylor saw that many areas of science were becoming increasingly reliant on new ways of collaborative, multidisciplinary working. The term e-science is intended to capture these new modes of working:

e-Science is about global collaboration in key areas of science and the next generation of infrastructure that will enable it.

(John Taylor, <http://www.e-science.clrc.ac.uk>)

Implicitly, these e-science applications define a set of computational and data services that the hardware and middleware infrastructure must deliver to enable this science revolution. The hardware will be of many forms: high-speed networks, super-

[†] Permanent address: Department of Electronics and Computer Science, University of Southampton, Southampton SO17 1BJ, UK.

One contribution of 13 to a Theme 'Information, knowledge and technology'.

computers, clusters of workstations and expensive shared experimental facilities. In this paper we shall concentrate almost entirely on the functionality of the middleware that is required for this e-science vision to become a reality. Despite the manifest risk of confusion, we shall refer to this middleware infrastructure supporting e-science as the ‘Grid’ (Foster & Kesselman 1999), although we recognize that the capabilities of current Grid middleware fall far short of these requirements! It is important to note that this definition of the Grid is very different from many public perceptions of ‘the Grid’—these range from connecting supercomputers to mass harvesting of idle cycles on millions of personal computers, as in the SETI@home project. Nor is our version of Grid middleware concerned only with ‘big science’, such as particle physics and astronomy. Much more relevant is the vision of Foster and co-workers of the Grid as providing the middleware infrastructure to enable construction of ‘virtual organizations’ (Foster *et al.* 2001). In this context, the Grid has clear relevance not only to the scientific community but also to industry and commerce. Our emphasis is thus on Grid middleware that enables dynamic interoperability and virtualization of IT systems, rather than Grid middleware to connect high-performance computing systems, to exploit idle computing cycles or to do ‘big science’ applications.

The plan of the paper is as follows. In the next section we begin with a brief discussion of the technological context of e-science and the Grid. We then give a brief description of the US National Aeronautics and Space Administration (NASA) Information Power Grid (IPG) as an early example of a ‘prototype production Grid’ (<http://www.nas.nasa.gov/About/IPG/ipg.html>). Our discussion of e-science and the Grid is set in the context of the UK e-Science Programme. After a brief overview of this programme we highlight some UK e-science projects in science, engineering and medicine (Hey & Trefethen 2002). The next section looks at the open-standards approach to the Grid being adopted by the Global Grid Forum and compares this with community-based standardization processes for the Internet, message-passing interface (MPI), Linux and the Web. Section 6 looks at some implications of the imminent data deluge that will arise from the new generation of e-science experiments in terms of archiving and curation. We conclude with some remarks about both social and technological issues posed by Grid-enabled ‘collaboratories’ in both scientific and commercial contexts.

2. Technology drivers for e-science and grids

The two key technological drivers of the IT revolution are Moore’s law (the exponential increase in computing power and solid-state memory) and the dramatic increase in communication bandwidth made possible by optical-fibre networks using optical amplifiers and wave-division multiplexing. In a very real sense, the actual cost of any given amount of computation and/or sending a given amount of data is falling to zero. Needless to say, while this statement is true for any fixed amount of computation and for the transmission of any fixed amount of data, scientists are now attempting calculations requiring orders of magnitude more computing and communication than was possible only a few years ago. Moreover, in many future experiments they are planning to generate several orders of magnitude more data than has been collected in the whole of human history.

The highest performance supercomputing systems of today consist of several thousands of processors interconnected by a special-purpose, high-speed, low-latency network. On appropriate problems it is now possible to achieve sustained performance of tens of teraflops (a million million floating-point operations per second). In addition, there are experimental systems under construction aiming to reach petaflops speeds within the next few years (Allen *et al.* 2001; Sterling 2002). However, these very high-end systems are, and will remain, scarce resources located at a relatively small number of sites. The vast majority of computational problems do not require such expensive massively parallel processing and can be satisfied by the widespread deployment of cheap clusters of computers at university, department and research group level.

The situation for data is somewhat similar. There is a relatively small number of centres around the world that act as major repositories of a variety of scientific data. Bioinformatics, with its development of gene and protein archives, is an obvious example. The Sanger Centre at Hinxton near Cambridge, UK (<http://www.sanger.ac.uk>), currently hosts 20 terabytes of key genomic data and has a cumulative installed processing power (in clusters, not a single supercomputer) of many teraflops. The Sanger Institute estimates that the amount of genome sequence data is increasing by a factor of four each year and that the associated computer power required to analyse these data will ‘only’ increase by a factor of two each year—still significantly faster than Moore’s law.

A different data/computing paradigm is apparent for the particle physics and astronomy communities. In the next decade we will see new experimental facilities coming on-line that will generate datasets ranging in size from hundreds of terabytes to tens of petabytes per year. Such enormous volumes of data exceed the largest commercial databases currently available by one or two orders of magnitude (Gray & Hey 2001). Particle physicists are energetically assisting in building Grid middleware† that will allow them not only to distribute these data to 100 or more sites and to more than 1000 physicists collaborating in each experiment, but also to perform sophisticated distributed analysis, computation and visualization on both subsets and the totality of the data. Particle physicists envisage a data/computing model with a hierarchy of data centres with associated computing resources, distributed around the global collaboration.

From these and other examples given below it will be evident that the volume of e-science data generated from sensors, satellites, high-performance-computer simulations, high-throughput devices, scientific images and new experimental facilities will soon dwarf that of all of the scientific data collected in the whole history of scientific exploration. Moreover, until very recently, commercial databases have been the largest repositories for electronically archiving data for further analysis. These commercial archives usually make use of relational database management systems (RDMSs) such as Oracle, IBM’s DB2 or Microsoft’s SQLServer. Today, the largest commercial databases range from tens of terabytes up to a few hundred terabytes. In the near future, this situation will change dramatically and the volume of data in scientific data archives will vastly exceed that in commercial RDMSs.

† EU Data Grid Project, <http://www.eu-datagrid.web.cern.ch>; NSF GriPyNProject, <http://www.griphyn.org>; DOE PPDataGrid Project, <http://www.ppdg.net>; UK GridPP Project, <http://www.gridpp.ac.uk>; NSF international Virtual Data Grid Laboratory (iVDGL) Project, <http://www.ivdgl.org>.

Inevitably, this watershed brings with it challenges and opportunities. For this reason, we believe that the data access, integration and federation capabilities of the next generation of Grid middleware will play a key role both for e-science and for e-business.

3. The NASA IPG

Over the last four or five years, NASA has pioneered a new style of computing infrastructure by connecting the computing resources of several of its research and development (R&D) laboratories to form the IPG (Johnston 2003). The vision for the IPG has been most clearly enunciated by Bill Johnston, who led the IPG implementation activity. Johnston says that the IPG is intended to

promote a revolution in how NASA addresses large-scale science and engineering problems by providing *persistent infrastructure* for ‘highly capable’ computing and data management services that, on-demand, will locate and co-schedule the multi-Center resources needed to address large-scale and/or widely distributed problems, and provide the ancillary services that are needed to support the workflow management frameworks that coordinate the processes of distributed science and engineering problems.

In NASA’s view, such a framework is necessary for their organization to address the problem of simulating *whole* systems. It is not only computing resources but also expertise and know-how that are distributed. In order to simulate a whole aircraft (wings, engines, landing gear and so on) NASA must bring together not only the necessary computing resources but also support mechanisms allowing engineers and scientists located at different sites to collaborate. This is the ambitious goal of the IPG. In the context of our discussion below, the IPG can be classified as an ‘intra-grid’, i.e. a grid connecting different sites with different expertise and hardware resources but all within one organization or enterprise.

In his presentations on the IPG, Johnston also makes the point that, although technology capable of enabling such distributed computing experiments has been around for many years, experiments on large-scale distributed computing have tended to be ‘one off’ demonstration systems requiring the presence and knowledge of expert distributed system engineers. He argues persuasively that what we need is grid middleware that makes the construction of such systems *routine*, with no need for experts to build and maintain the system. The IPG is the first step along the road towards a ‘production-quality’ grid. NASA based their IPG middleware largely on Globus (<http://www.globus.org>), a toolkit that offers secure ‘single sign-on’ functionality to distributed resources through digital certificate technology (Foster & Kesselman 1997). This technology authenticates user access to multiple, distributed resources but only to those for which they have prior authorization. This eliminates the need for cumbersome multiple log-ins, passwords and so on. Users are then able to obtain data from multiple remote data repositories and schedule computation and analysis of these data on the most convenient computing resource to which they are allowed access. Moreover, NASA’s vision for the IPG incorporates the development of ‘workflow’ management software to support the steps that users need to perform to assemble complex data access requests and co-schedule the necessary computational

resources. These are their requirements for the grid middleware infrastructure that will empower NASA scientists to attack the next generation of complex distributed science and engineering problems.

4. The UK e-Science Programme

In the UK, John Taylor was able to back his vision for e-science by allocating £120M to establish a three-year e-science R&D programme (Hey & Trefethen 2002). This e-science initiative spans all the research councils: the Biotechnology and Biological Sciences Research Council, the Council for the Central Laboratory of the Research Councils (CCLRC), the Engineering and Physical Sciences Research Council, the Economic Social Research Council, the Medical Research Council, the Natural Environment Research Council and the Particle Physics and Astronomy Research Council. A specific allocation was made to each research council, with particle physicists being allocated the lion's share so that they could begin putting in place the middleware infrastructure necessary to exploit the CERN Large Hadron Collider (LHC) experiments that are projected to come on-stream in 2007, as mentioned earlier. All of the research councils have now selected their pilot e-science projects: the major goals of some of these projects are described below. The allocation to the Daresbury and Rutherford Laboratories by CCLRC is specifically to 'grid-enable' their experimental facilities. In addition, £10M was allocated towards the procurement of a new national teraflops computing system. The remaining £15M of the 2000 Government Spending Review (SR2000) e-science allocation was designated for an e-science 'Core Programme'. This sum was augmented by £20M from the Department of Trade and Industry for an industry-facing component of the Core Programme. This £20M requires a matching contribution from industry. So far, industrial contributions to the e-science programme exceed £30M and over 50 UK and US companies are involved in collaborative projects.

The goal of the Core Programme is to support the e-science pilot projects of the different research councils and to work with industry to accelerate the development of robust, 'industrial strength', generic grid middleware. Requirements and lessons learnt from the different e-science applications will inform the development of more stable and functional grid middleware that can both assist the e-science experiments and but also be of relevance to industry and commerce. The Core Programme has set up e-science centres around the country to form the basis for a UK e-science grid. Together with the involvement of the Daresbury and Rutherford Laboratories, the goal for this UK e-science grid is to develop a critical mass of expertise in building a production grid. With the experience gained from this prototype, the Core Programme can assist the e-science projects in building their discipline-centric grids and learn how to interconnect and federate multiple grids in a controlled way. If we can build such grids to production quality, operating 24 hours a day, seven days a week, 52 weeks a year, and that do not require computing experts to use them, then e-science really does have the potential to change the way we do scientific research in all our universities and research institutes.

We now briefly describe the goals of some of the UK e-science projects to illustrate what the new generation of e-scientists wish to be able to do and what functionality the grid middleware must support to enable them to carry out their science. These e-

science applications also illustrate the spectacular growth forecast for scientific data generation.

(a) *e-Chemistry: the Combechem project*

The Combechem project (<http://www.combechem.org>) is concerned with the synthesis of new compounds by combinatorial methods and mapping their structure and properties. It is a collaboration between the Universities of Southampton and Bristol, the Cambridge Crystallographic Data Centre and the companies Pfizer and IBM. Combinatorial methods provide new opportunities for the generation of large amounts of new chemical knowledge. Such a parallel synthetic approach can create hundreds of thousands of new compounds at a time and this will lead to an explosive growth in the volume of data generated. These new compounds need to be screened for their potential usefulness and their properties and structure need to be identified and recorded. An extensive range of primary data therefore needs to be accumulated and integrated, and relationships and properties modelled. The goal of the Combechem project is therefore to develop an integrated platform that combines existing structure and property data sources within a grid-based information and knowledge sharing environment.

The first requirement for the Combechem platform is support for the collection of new data, including process as well as product data, based on integration with electronic laboratory and e-logbook facilities. The next step is to integrate data generation on demand via grid-based quantum and simulation modelling to augment the experimental data. In order for the environment to be usable by the community at large, it will be necessary to develop interfaces that provide a unified view of resources, with transparent access to data retrieval, online modelling and design of experiments to populate new regions of scientific interest.

This service-based grid-computing infrastructure will extend to devices in the laboratory as well as to databases and computational resources. In particular, an important component of the project is the support of remote users of the National Crystallographic Service, which is physically located in Southampton, UK. This support extends not only to portal access but also to support for the resulting workflow—incorporating use of the X-ray e-laboratory, access to structures databases and computing facilities for simulation and analysis. The resulting environment is designed to provide shared, secure access to all of these resources in a supportive collaborative e-science environment.

(b) *Bioinformatics: the myGrid project*

The myGrid project (<http://mygrid.man.ac.uk>) is led by the University of Manchester and is a consortium comprising the Universities of Manchester, Southampton, Nottingham, Newcastle and Sheffield together with the European Bioinformatics Institute and industrial partners GSK, AstraZeneca, IBM and SUN. The goal of myGrid is to design, develop and demonstrate higher-level functionality grid middleware to support scientists' use of complex distributed resources. An e-scientist's workbench, similar in concept to that of the Combechem project described above, will be developed to support the scientific process of experimental investigation, evidence accumulation and result assimilation. It will also help the scientist's use of

community information and enhance scientific collaboration by assisting the formation of dynamic groupings to tackle emergent research problems.

A novel aspect of the workbench is its support for individual scientists by provision of personalization facilities relating to resource selection, data management and process enactment. The design and development activity will be informed by and evaluated using two specific problems in bioinformatics. The bioinformatics community is typically highly distributed and makes use of many shared tools and data resources. The myGrid project will develop two application environments, one that supports the analysis of functional genomic data, and another that supports the annotation of a pattern database. Both of these tasks require explicit representation and enactment of scientific processes, and have challenging performance requirements.

It is perhaps useful to give some idea of the types and volumes of data involved in such bioinformatics problems. The European Bioinformatics Institute in the UK is one of three primary sites in the world for the deposition of nucleotide sequence data. It currently contains around 14×10^6 entries of 15×10^9 bases with a new entry received every 10 seconds. Data at the three centres (in the USA, UK and Japan) is synchronized every 24 hours. This database has tripled in size in the last 11 months. About 50% of the data is for human DNA, 15% for mouse DNA and the rest for a mixture of organisms. The total size of the database is of the order of several terabytes. As a second example, we consider gene expression databases that involve image data produced from DNA chips and microarrays. In the next few years we are likely to see hundreds of experiments in thousands of laboratories world wide and the consequent data-storage requirements are predicted to be in the range of petabytes per year. Lastly, consider two of the most widely used protein databases: the Protein Data Bank (PDB) and SWISS-PROT. The PDB is a database of three-dimensional protein structures. At present there are around 20 000 entries and around 2000 new structures are being added every 12 months. The total database is quite small, of the order of gigabytes. SWISS-PROT is a protein sequence database currently containing around 100 000 different sequences with knowledge abstracted from around 100 000 different scientific articles. The present size is of the order of tens of gigabytes with an 18% increase over the last eight months. The challenge for the bioinformatics community is to extract useful biological or pharmaceutical knowledge from this welter of disparate sources of information. The myGrid project, with its goal of providing tools for automatic annotation and provenance tracking, aims to take some first steps along this road from data to information to knowledge.

(c) *e-Engineering: the DAME project*

An important engineering problem for industry is concerned with health monitoring of industrial equipment. The Distributed Aircraft Maintenance Environment (DAME, <http://www.cs.york.ac.uk/DAME>) project is a consortium of university research groups (from York, Leeds, Sheffield and Oxford) with the participation of BAe Systems and Rolls-Royce. The problem is the analysis of sensor data generated by the many thousands of Rolls-Royce engines currently in service. For example, each transatlantic flight made by each engine generates about a gigabyte of data per engine: from pressure, temperature and vibration sensors. The goal of the project is to transmit a small subset of these primary data for analysis and comparison with engine data stored in one of several data centres located around the world.

By identifying the early onset of problems, Rolls-Royce hope to be able to lengthen the period between scheduled maintenance periods, thus increasing profitability. The engine sensors will generate many petabytes of data per year and decisions need to be taken in real-time as to how many data to analyse, how many to transmit for further analysis and how many to archive. Similar (or larger) data volumes will be generated by other high-throughput sensor experiments in fields as varied as environmental and earth observation, as well as, of course, in human health-care monitoring.

(d) *e-engineering: the GEODISE project*

The Grid Enabled Optimization and Design Search for Engineering (GEODISE, <http://www.geodise.org>) project is a collaboration between the Universities of Southampton, Oxford and Manchester, together with the companies BAe Systems, Rolls-Royce and Fluent. The project aims to provide grid-based seamless access to a state-of-the-art collection of optimization and search tools, industrial strength analysis codes, distributed computing and data resources and an intelligent knowledge repository. Besides traditional engineering design tools such as computer-aided design systems, computational fluid dynamics (CFD) and finite-element-method (FEM) simulations on high-performance clusters, multi-dimensional optimization methods and interactive visualization techniques, the project is working with engineers at Rolls-Royce and BAe Systems to capture knowledge learnt in previous product design cycles.

Engineering design search and optimization is the process whereby engineering modelling and analysis are exploited to yield improved designs. In the next two to five years, intelligent search tools will become a vital component of all engineering design systems and will steer the user through the process of setting up, executing and post-processing design search and optimization activities. Such systems typically require large-scale distributed simulations to be coupled with tools to describe and modify designs using information from a knowledge base. These tools are usually physically distributed and under the control of multiple elements in the supply chain. While evaluation of a single design may require the analysis of gigabytes of data, to improve the process of design can require the integration and interrogation of terabytes of distributed data. Achieving the latter goal will lead to the development of intelligent search tools. The application area of focus is that of CFD, which has clear relevance to the industrial partners.

(e) *e-Healthcare: the eDiamond project*

The Grid will open up enormous opportunities in the area of collaborative working at widely distributed sites, and, in particular, to the effective sharing of large, federated databases of images. The eDiamond (<http://www.gridoutreach.org.uk/docs/pilots/ediamond.htm>) project aims to exploit these opportunities to create a world-class resource for the UK mammography community. Medical-image databases represent both huge challenges and huge opportunities. Medical images tend to be large, variable across populations, contain subtle clinical signs, have a requirement for loss-less compression, require extremely fast access, have variable quality, and have privacy as a major concern. The variable quality of mammograms and the scarcity of breast-specialist radiologists are major challenges facing healthcare systems here and throughout the world.

The eDiamond project brings together medical-image analysis expertise from Mirada Solutions Ltd (a UK start-up company), the MIAS Interdisciplinary Research Collaboration (Medical Images and Signals to Clinical Information), which is already heavily involved in developing Grid applications, computer science expertise from IBM and the Oxford e-science Centre, and clinical expertise from the Scottish Breast Screening Centres (Edinburgh and Glasgow), the Oxford Radcliffe Trust and St George's, Guy's and St Thomas' NHS Trust hospitals in London. eDiamond aims to provide an exemplar of the dynamic, best-evidence-based approach to diagnosis and treatment made possible through the Grid. The Oxford e-Science Centre is working with IBM to create a 'virtual organization', comprising clinical sites at the Churchill Hospital in Oxford, St George's and Guy's Hospitals in London and Breast Screening Centres in Scotland.

This project requires that several generic e-science challenges be addressed both by leveraging existing grid technology and by developing novel middleware solutions. Key issues include the development of each of the following:

- (i) ontologies and metadata for the description of demographic data, the physics underpinning the imaging process, key features within images and relevant clinical information;
- (ii) large, federated databases both of metadata and images;
- (iii) data compression and transfer;
- (iv) effective ways of combining grid-enabled databases of information that must be protected and that are based in hospitals protected by firewalls;
- (v) very rapid data-mining techniques
- (vi) a secure grid infrastructure for use within a clinical environment.

A large federated database of annotated mammograms will be built up for which all mammograms entering the database will be standardized prior to storage to ensure database consistency and reliable image processing. Such a database provides the basis for new applications in teaching, aiding detection, and aiding diagnosis that will be developed as part of the project.

(f) *Particle physics: GridPP and the LHC Grid*

The world-wide particle physics community is planning an exciting new series of experiments to be carried out by the new LHC experimental facility under construction at CERN in Geneva. The goal is to find signs of the Higgs boson, key to the generation of mass for both the vector bosons and the fermions of the Standard Model of the weak and electromagnetic interactions. Particle physicists are also hoping for indications of other new types of matter (such as supersymmetric particles) which may shed light on the 'dark matter' problem of cosmology. As noted earlier, these LHC experiments are on a scale never before seen in physics, with each experiment involving a collaboration of over 100 institutions and over 1000 physicists from Europe, USA and Japan. When operational in 2007, the LHC will generate petabytes of experimental data per year, for each experiment. This vast amount of data needs to be pre-processed and distributed for further analysis by all

members of the consortia to search for signals betraying the presence of the Higgs boson or other surprises. The physicists need to put in place an LHC grid infrastructure that will permit the transport and data mining of extremely large and distributed datasets. The GridPP project (<http://www.gridpp.ac.uk>) is a consortium of UK particle physics research groups who are building the UK component of the world-wide infrastructure for the LHC grid. The GridPP consortium also plays a key role in the EU-funded DataGrid Project (EU DataGrid Project, <http://www.eu-datagrid.web.cern.ch>). Both these projects liaise closely with the CERN LHC Grid development team and the three major particle physics projects in the USA (the NSF GriPhyN project, <http://www.griphyn.org>; the UK DOE GridPP Project, <http://www.gridpp.ac.uk>; and the NSF iVDGL project, <http://www.ivdgl.org>). A complementary project, the EU-funded DataTAG (<http://www.datatag.org>) project, is exploring end-to-end quality of service issues for transatlantic networks.

The LHC at CERN in Geneva will begin to generate collision data in early 2007. Two of the initial four experiments at the LHC are the ATLAS and CMS global collaborations. Each of these LHC experiments will need to store, access and process *ca.* 10 petabytes per year and will require the use of some 200 teraflops of processing power for reconstruction, analysis and simulation. By 2015, particle physicists will be using exabytes of storage and petaflops of (non-supercomputer) computation. At least initially, it is likely that most of these data will be stored in a distributed file system, with the associated metadata stored in a relational database. CERN cannot provide all the computing and storage facilities required by these experiments, so use of grid middleware to share computation and data across the collaborations is seen as essential. The particle physicists clearly have extreme demands for data analysis and storage, but these data are all of the same general character: the records of independent collision events. Since each event may be treated independently of any other, high-capacity cluster-type computing rather than high-capability supercomputer time is sufficient for their reconstruction, analysis and simulation needs. Nonetheless, the LHC grid will have to be able to handle the distribution and storage of huge amounts of such data.

(g) Astronomy: the AstroGrid project and the International Virtual Observatory

This e-science project is very much data-centric. In the UK, the astronomers are planning to create a ‘virtual observatory’ in the AstroGrid project (<http://www.astrogrid.ac.uk>). There are similar initiatives in the USA with the NSF NVO project (<http://www.nvo.org>) and in Europe with the EU AVO project (<http://www.eso.org/avo>). The goal of these projects is to provide uniform access to a federated, distributed repository of astronomical data, spanning all wavelengths from radio waves to X-rays. At present, astronomical data using different wavelengths are captured by different telescopes and stored in a variety of formats. The goal is to create a ‘data warehouse’ for astronomical data that will enable new types of studies to be performed. Astronomers in Europe and the USA are working together to build a suitable grid infrastructure to support these virtual observatories.

At present, the largest astronomy database is *ca.* 10 terabytes. However, new telescopes coming online will radically change this picture. For example, it is estimated that the NVO project alone will store 500 terabytes per year from 2004. Similarly, the Laser Interferometer Gravitational Observatory (LIGO) project is estimated to

generate 250 terabytes per year beginning in 2002 (<http://www.ligo.caltech.edu>). A final example of how the astronomical landscape will be revolutionized in the next decade is the VISTA survey project (<http://www.vista.ac.uk>) in the visible and infrared regions. The VISTA telescope will be operational from 2004 and will generate 250 gigabytes of raw data per night and *ca.* 10 terabytes of stored data per year. In 10 years or so, there will be several petabytes of data in the VISTA archive.

In order for these world-wide virtual observatory efforts to be successful, the astronomy community needs to work together to agree on metadata standards to describe astronomical data. Members of AstroGrid, NVO and the AVO projects met in June 2002 and formed the International Virtual Observatory Alliance (<http://www.ivoa.net>). As a first step, this body has now agreed on a standard 'VOTable' format for astronomical results. In addition it is recognized that there will need to be compatibility with the widely used Flexible Image Transport System (FITS) standard. It is also likely that we will see the emergence of an Astronomical Query Language 'AQL', specified in XML. The existence of such standards for metadata will be vital for the interoperability and federation of astronomical data held in different formats in file systems, databases or other archival systems. To drive and test the middleware development, the AstroGrid project has also identified a list of the 'top ten' science problems for the virtual observatory.

5. Open-standard grid middleware

In the UK context, we regard the IPG as an 'existence proof' that the grid vision for e-science and e-engineering can be realized. Nevertheless, it is clear from the descriptions of the e-science experiments above that their grid middleware requirements are many and varied. Single sign-on with secure authentication as provided by Globus digital certificates is a useful starting point, but we need to work with the world-wide grid community to enhance the functionality and reliability of the present incarnation of the grid middleware.

The vision for a layer of 'Grid' middleware that provides a set of core services to enable such new types of science and engineering is due to Ian Foster, Carl Kesselman and Stephen Tuecke (Foster *et al.* 2001). Within the Globus project, they have developed parts of a prototype open-source Grid Toolkit (Foster & Kesselman 1997). Their choice of the name 'grid' to describe this middleware infrastructure resonates with the idea of a future in which computing resources, compute cycles and storage, as well as expensive scientific facilities and software, can be accessed on demand like the electric power utilities of today. These 'e-utility' ideas are also reminiscent of the recent trend of the Web community towards a model of 'Web services' advertised by brokers and consumed by applications. The recent move of the grid community towards a service-oriented architecture and the proposal for an Open-Grid-Services Architecture (OGSA) based on commercially supported Web services technology is therefore of great significance. The UK community is participating in international standardization efforts via the Global Grid Forum (<http://www.globalgridforum.org>). The Globus toolkit will develop into an open-source version of enhanced grid middleware conforming to the new OGSA standard (Foster *et al.* 2003). We also see it as beneficial that there are likely to be multiple commercial implementations of the OGSA-compliant middleware.

It is interesting to note that the Global Grid Forum is a community-based standardization process along the lines of the Internet Engineering Task Force (IETF), the Message Passing Interface (MPI) Forum and the World Wide Web Consortium (W3C). Rather than go through lengthy and often contentious official standardization processes under the auspices of official standards bodies, in these pioneering efforts researchers and industry have come together to hammer out an ‘informal’ standard that has wide support across their communities. In the case of the MPI standard, a series of working meetings converged on the standard within a year. The availability of an open-source version of this emerging MPI standard, produced by Rusty Lusk and Bill Gropp from Argonne National Laboratory, was crucial to the success of this process (Hempel & Walker 1999). There are now both open-source and commercial implementations of MPI available. These examples, together with the example of the Linux community, demonstrate the power of such a community-based approach to standards. It is fair to say that the Internet, Linux and the Web have changed the course of major global corporations and have affected much of the IT industry. It is clear that the Grid has the potential to change the future of e-business as profoundly as the Internet, Linux and the Web.

6. Scientific metadata, information and knowledge

Metadata are data about data. We are all familiar with metadata in the form of catalogues, indices and directories. Librarians work with books that have a metadata ‘schema’ containing information such as title, author, publisher and date of publication at the minimum. On the World Wide Web, most Web pages are coded in the HyperText Mark-up Language (HTML). This contains instructions as to the appearance of the page (size of headings and so on) as well as hyperlinks to other Web pages. Recently, the eXtensible Mark-up Language (XML) has been agreed by the W3C standards body. XML allows Web pages and other documents to be tagged with computer-readable metadata. The XML tags give information about the structure of the data contained in the document rather than instructions as to presentation. For example, XML tags could be used to give an electronic version of the book schema given above.

The quality of the metadata describing the data is important. Search engines to extract meaningful information can be constructed from the metadata that are annotated in documents stored in electronic form. The quality of the search engine will only be as good as the metadata that it references. There is now a movement to standardize other, ‘higher-level’ mark-up languages, such as DARPA Agent Markup Language (DAML) and Ontology Inference Language (OIL) (<http://www.daml.org/>), which would allow computers to be able to reason about the ‘meaning or semantic relationships’ contained in a document. This is the ambitious goal of Tim Berners-Lee’s ‘semantic Web’ (Berners-Lee *et al.* 2001).

Although we have given a simple example of metadata in relation to textual information, metadata will also be vital for storing and preserving scientific data. Such scientific data metadata will not only contain information about the annotation of data by semantic tags, but will also provide information about its provenance and its associated user access controls. In order to construct ‘intelligent’ search engines, each separate community and discipline needs to come together to define generally accepted metadata standards for their community data grids. Furthermore, just as

the Web is attempting to move beyond information to knowledge, so the same scientific communities will need to define relevant ‘ontologies’—roughly speaking, relationships between the terms used in shared and well-defined vocabularies for their fields—that can allow the construction of ‘semantic grids’ (DeRoure *et al.* 2003; Moore 2001).

With the imminent data deluge, the issue of how we handle this vast outpouring of scientific data becomes of paramount importance. Up to now, we have generally been able to manually manage the process of examining the experimental data to identify potentially interesting features and discover significant relationships between them. In the future, when we consider the massive amounts of data being created by simulations, experiments and sensors, it is clear that in many fields we will no longer have this luxury. We therefore need to automate the discovery process, from data to information to knowledge, as far as possible. At the lowest level, this requires automation of data management with the storage and organization of digital entities. Next we need to move towards automatic information management. This will require automatic annotation of scientific data with metadata describing interesting features both of the data and of its storage and organization. Finally, we need to attempt to progress beyond structure information towards automated knowledge management of our scientific data. This will include the expression of relationships between information tags as well as information about the storage and organization of such relationships.

There is one last issue we wish to highlight. In the future, we envisage that scientific data, whether generated by direct experimental observation or by *in silico* simulations on supercomputers or clusters, will be stored in a variety of ‘data grids’. Such data grids will involve data repositories together with the necessary computational resources required for analysis, distributed around the global e-science community. The scientific data, held in file stores, databases or archival systems, together with a metadata catalogue, probably held in an industry standard relational database, will become a new type of distributed and federated digital library. Up to now the digital-library community has been primarily concerned with the storage of text, audio and video data. The scientific digital libraries that are being created by global, collaborative e-science experiments will need the same sort of facilities as conventional digital libraries: a set of services for manipulation, management, discovery and presentation. In addition, these scientific digital libraries will require new types of tools for data transformation, visualization and data mining. The community will also need to solve the problem of the long-term curation of such data and its ancillary data-manipulation programs.

Generating the data is one thing, preserving them in a form so that they can be used by scientists other than the creators is entirely another issue. This is the process of ‘curation’. For example, the SWISS-PROT (<http://www.ebi.ac.uk/swissprot>) database is generally regarded as the ‘gold standard’ for protein structure information. For SWISS-PROT, curation is done by a team of 25 full-time curators split between the Swiss Bioinformatics Institute and the European Bioinformatics Institute (EBI). This illustrates the present labour-intensive nature of the curation process and why it will be necessary to address this support issue, involving a mix of automated, semi-automated and manual annotation and data cleansing. In addition, long-term preservation of the data, and of the software environment that produced the data, will be a crucial aspect of the work of a data repository. There are many

technical challenges to be solved to ensure that the information generated today can survive changes in storage media, devices and digital formats (Rothenberg 1995). Needless to say, a solution to these problems is much more than just a technical challenge: all parts of the community from digital librarians and scientists to computer scientists and IT companies will need to be involved.

7. Conclusions

There are many interesting technical, social and legal issues that arise from the above vision of e-science and the Grid. Here we briefly mention some of them.

The grid standardization activity via the Global Grid Forum is a community activity that brings together application scientists, computer scientists, IT specialists and the IT industry. Clearly, such a multi-disciplinary endeavour will inevitably create tensions, for example, between academia and industry, between research and development. Such tension (between research and development) is also evident in the UK e-science application projects. At an extreme we could approach a situation in which the application scientists just want to do their application research, the computer scientists want to do their computer science research and no one wants to write and document robust re-usable grid middleware! Keeping all parties interested and feeling that they gain more by collaboration and doing science in new ways is a delicate balancing act. Alleviating this problem is one rationale for the part of the UK's Core Programme looking at generic requirements of the e-science projects. By working with commercial software companies we can hope to develop robust middleware solutions that are effective as well as being documented and maintainable.

Another interesting potential problem posed by the emergence of e-science 'collaboratories' is the provenance not only of data but of intellectual property (IP). When multiple organizations are collaborating to create new IP, it may be crucial to be able to identify who invented what, when. The myGrid project is a good example in which several universities, two IT companies and three pharmaceutical companies are engaged in the virtual organization. For this reason the project is developing ways to record and track the provenance of both data and IP. This area is likely to be the focus of a major thrust for the UK Programme, and for industry.

Security has many aspects, and here we mention only two such 'security' issues. The first is the requirement for IT managers in organizations to revise their traditional approach to security. Usually, IT managers will install a system of firewalls, rather like a series of moats round a castle, so that breaching of one firewall by an unauthorized intruder does not compromise the entire enterprise. In the case of grids and virtual organizations, these same managers have to be convinced that they can trust the digital certificate authentication and allow users not registered at their site to access 'their' resources. A different security issue arises in the context of collaborations involving medical data. Here the issue is not so much about firewalls and intruders as about privacy and anonymization of patient data. This is a critical issue that need to be resolved to the satisfaction of all stakeholders before the grid is to make significant inroads in delivering on the promise of e-healthcare.

Many social and legal issues are bound also to arise due to the need for trust relationships within a virtual organization. In some of the scenarios described above, where intelligent agents or services are at work on behalf of a grid user, issues of

contractual obligation may come into play if and when an effective or optimal service is not provided. With virtual organizations trust policies become key in such automatic negotiations and commitment.

The Virtual Observatory projects raise interesting social issues that have relevance to more than just the astronomy community. Increasingly, in many countries there is an emerging consensus that researchers have a duty to make primary research data obtained from projects funded from public funds available to the wider scientific community (Wouters 2002). This is often a very controversial issue (see also Nelson 2003). Researchers, from many disciplines, who have painstakingly accumulated what they regard as valuable research data, have a natural reluctance to allow other researchers to mine ‘their’ data and obtain research rewards. Many ‘defensive’ reasons are often given to ‘explain’ why it is inappropriate or meaningless to place the primary data in the public domain. Certainly, in order to be useful, the primary data need to be annotated by high quality metadata describing the particulars of the experimental instrument, methodology, and so on. Who will do this annotation? Unless funding agencies make the creation of such metadata and the placement of data generated by the research project into the public domain a condition of the grant, it is likely that any such effort will founder. And, of course, there will have to be exceptions to allow legitimate protection and exploitation of IP created in research projects. There is also a real cost to this long-term curation of data: many scientists will see this as money taken away from their budget to do ‘real’ research. Nevertheless, if something like the Virtual Observatory succeeds it will assist in the ‘democratization’ of astronomy. Scientists who were unsuccessful in being awarded precious observing time at the small number of large telescopes around the world will be able to do their research on data obtained by others. Who gets the credit? Will users be obliged to cite the provenance of the data they have used? It has been argued that alongside the two traditional methodologies of science—theory and experiment—computational science has now emerged as a third methodology. With the advent of scientific data warehouses such as virtual observatories we may be seeing the emergence of a fourth methodology, that of collection-based science.

To sum up: the e-science vision is extremely ambitious. It dares to envisage the creation of a new and powerful middleware infrastructure that supports a whole new way of doing collaborative science. If successful, it will enable different communities to come together and create robust, secure, virtual organizations to attack new and complex problems, exploiting a wide variety of distributed resources. In the context of virtual organizations and e-utility models of resources-on-demand, the Grid also promises to be a significant development for industry. Although inspired by the demands of distributed scientific applications, the Grid, like the Internet, Linux and the Web before it, has the potential to be at least as disruptive a technology for industry as these earlier developments that also had their origins in academia. The scientific, engineering and medical applications offer a very demanding set of requirements for grid middleware and these applications will provide a stringent testing ground for grid infrastructure. The imminent flood of data from the new generation of high-throughput technologies applied to many fields will also pose significant challenges for computer science to assist the process of going from data to information and finally to knowledge. Similarly, IBM’s vision (http://www.research.ibm.com/autonomic/manifesto/agenda2001_p1.html) for ‘autonomic computing’—as in a layer of grid middleware with autonomic capabilities for self-configuring, self-optimization,

self-healing and self-protecting—poses significant research challenges for computer science. The e-science/Grid vision is nothing if not ambitious. We have made a start and are past ‘base camp’, but there is a long road ahead!

Appendix A. Web addresses for projects mentioned in the paper

All addresses use the hypertext transfer protocol (<http://>).

AstroGrid Project	www.astrogrid.ac.uk
Astrophysical Virtual Observatory (AVO)	www.euro-vo.org/intro.html
Combechem project	www.combechem.org
DAML and OIL (high-level mark-up languages)	www.daml.org
DAME Project	www.cs.york.ac.uk/DAME
DataTag Project	www.datatag.org
DOE PPDataGrid Project	www.ppdg.net
eDiamond Project	www.gridoutreach.org.uk/docs/pilots/ediamond.htm
EU DataGrid Project	eu-datagrid.web.cern.ch
GEODISE Project	www.geodise.org
Global Grid Forum	www.gridforum.org
Globus: a metacomputing infrastructure toolkit	www.globus.org
Gray & Hey (2001)	www.research.microsoft.com/~gray
P. Horn, <i>The state of information technology</i>	www.research.ibm.com/autonomic/manifesto/agenda2001_p1.html
Laser Interferometer Gravitational Wave Observatory (LIGO)	www.ligo.caltech.edu
LHC Computing Project	lhcgird.web.cern.ch/LHCgrid
myGrid Project	www.mygrid.man.ac.uk
NASA Information Power Grid	www.nas.nasa.gov/About/IPG/ipg.html
NSF GriPyN Project	www.griphyn.org
NSF iVDGL Project	www.ivdgl.org
NSF National Virtual Observatory	www.us-vo.org
Sanger Institute, Hinxton, UK	www.sanger.ac.uk
SWISS-PROT database	www.ebi.ac.uk/swissprot
J. M. Taylor	www.e-science.clrc.ac.uk
UK GridPP Project	www.gridpp.ac.uk

Visible & Infrared Survey
Telescope for Astronomy (VISTA)

www.vista.ac.uk

Wouters (2002)

dataaccess.ucsd.edu/NIWInew.htm

References

- Allen, F. (and 50 others) 2001 BlueGene: a vision for protein science using a petaflop computer. *IBM Syst. J.* **40**, 310.
- Berners-Lee, T., Hendler, J. & Lassila, O. 2001 The semantic Web. *Scient. Am.* **284**, 34–43.
- DeRoure, D., Jennings, N. & Shadbolt, N. 2003 The semantic grid: a future e-Science infrastructure. In *Grid computing: making the global infrastructure a reality* (ed. F. Berman, G. Fox & T. Hey), pp. 437–470. Wiley.
- Foster, I. & Kesselman, C. 1997 Globus: a metacomputing infrastructure toolkit. *Int. J. Supercomput. Applic.* **11**, 115–128.
- Foster, I. & Kesselman, C. (eds) 1999 *The grid: blueprint for a new computing infrastructure*. San Francisco, CA: Morgan Kaufmann.
- Foster, I., Kesselman, C. & Tuecke, S. 2001 The anatomy of the grid: enabling scalable virtual organizations. *Int. J. High-Performance Comput. Applic.* **15**, 200–222.
- Foster, I., Kesselman, C., Nick, J. & Tuecke, S. 2003 The physiology of the grid. In *Grid computing: making the global infrastructure a reality* (ed. F. Berman, G. Fox & T. Hey), pp. 217–250. Wiley.
- Gray, J. & Hey, T. 2001 In search of petabyte databases. In *Proc. 2001 HPTS Workshop, Asilomar, CA*. (Available at www.research.microsoft.com/~gray.)
- Hempel, R. & Walker, D. W. 1999 The emergence of the MPI message passing standard for parallel computing. *Comput. Stand. Interfaces* **7**, 51–62.
- Hey, T. & Trefethen, A. E. 2002 The UK e-Science Core Programme and the Grid. *Future Gener. Comput. Syst.* **18**, 1017–1031.
- Johnston, W. E. 2003 Implementing production grids. In *Grid computing: making the global infrastructure a reality* (ed. F. Berman, G. Fox & T. Hey), pp. 117–170. Wiley.
- Moore, R. W. 2001 Knowledge-based grids. *Proc. 18th IEEE Symp. Mass Storage Systems and 9th Goddard Conf. on Mass Storage Systems and Technologies, San Diego*. (CD ROM.)
- Nelson, R. R. 2003 The advance of technology and the scientific commons. *Phil. Trans. R. Soc. Lond. A* **361**, 1691–1708.
- Rothenberg, J. 1995 Ensuring the longevity of digital documents. *Scient. Am.* **272**, 24–29.
- Sterling, T. 2002 The Gilgamesh MIND processor-in-memory architecture for petaflops-scale computing. *ISHPC Conf., Kansai, Japan*, pp. 1–5.
- Wouters, P. 2002 Data sharing policies. OECD Report. (Available at <http://dataaccess.ucsd.edu/NIWInew.htm>.)

