

GATECloud.net: a Platform for Large-Scale, Open-Source Text Processing on the Cloud

Valentin Tablan, Ian Roberts, Hamish Cunningham, Kalina Bontcheva

*University of Sheffield
Department of Computer Science
Regent Court, 211 Portobello, S1 4DP
Sheffield, United Kingdom*

Cloud computing is increasingly being regarded as a key enabler of the “democratisation of science”, since on-demand, highly scalable cloud computing facilities enable researchers anywhere to carry out data-intensive experiments. In the context of Natural Language Processing (NLP), algorithms tend to be complex, which makes their parallelisation and deployment on cloud platforms a non-trivial task. This paper presents a new, unique cloud-based platform for large-scale NLP research – GATECloud.net. It enables researchers to carry out data-intensive NLP experiments by harnessing the vast, on-demand compute power of the Amazon cloud. Important infrastructural issues are dealt with by the platform, completely transparently for the researcher: load balancing, efficient data upload and storage, deployment on the virtual machines, security, and fault tolerance. We also include a cost-benefit analysis and usage evaluation.

1. Introduction

The continued growth of unstructured content and the availability of ever more powerful computers has resulted in an increased need for researchers in diverse fields (e.g. humanities, social sciences, bio-informatics) to carry out language processing and text mining experiments on very large document collections (or *corpora*). An additional impetus is the availability of key datasets, e.g. Wikipedia and Freebase snapshots, which can help with experimental repeatability. Many of these datasets are impossible to process in reasonable time on standard computers such as desktop machines or individual servers.

In the context of Natural Language Processing (NLP) research, large-scale algorithms (also referred to as data-intensive or web-scale NLP) are demonstrating increasingly superior results compared to approaches trained on smaller datasets, mostly thanks to addressing the data sparseness issue through collection of significantly larger numbers of naturally occurring linguistic examples [8]. The need for and the success of data-driven NLP methods to a large extent mirrors recent trends in other research fields, leading to what is being referred to as “the fourth paradigm of science” [3].

However, while researchers at big corporations (e.g. Google, Yahoo, Microsoft, IBM) have access to both Web-scale textual data (including web query logs) and vast computing infrastructures, scientists from smaller research groups and universities are faced with major technological challenges when carrying out cutting edge data-driven text processing experiments. Cloud computing [5] is increasingly being regarded

as a key enabler of the “democratisation of science” [7, 3], giving researchers everywhere affordable access to computing infrastructures, which allow the deployment of significant compute power on an on-demand basis, and with no upfront costs.

However, NLP algorithms tend to be complex, which makes deployment on cloud platforms a specialised, non-trivial task, with its own associated costs in terms of significant time overhead and expertise required.

To answer these challenges, we have developed a novel, unique cloud-based platform for large-scale NLP research – GATECloud.net. It aims to give researchers access to specialised software and enables them to carry out large-scale NLP experiments by harnessing the vast, on-demand compute power of the Amazon cloud. It also eliminates the need to implement specialised parallelisable text processing algorithms. Important infrastructural issues are dealt with by the platform, completely transparently for the researcher: load balancing, efficient data upload and storage, deployment on the virtual machines, security, and fault tolerance.

This paper is structured as follows. Section 2 differentiates GATECloud.net from related work on data-intensive NLP and cloud computing. Section 3 motivates the need for an NLP Platform-as-a-Service (PaaS) and presents a number of requirements. Next, the architecture and implementation of GATECloud.net are discussed (Section 4) in relation to these requirements. The paper concludes with a number of use cases and evaluation experiments (Section 5) and a discussion of future work (Section 6).

2. Large-Scale Text Mining and Compute Clouds

Following the Software-as-a-Service paradigm from cloud computing [5], a number of text processing services have been developed, e.g. OpenCalais¹, Extractiv², and Alchemy API³. These mostly provide information extraction services, which are accessible programmatically (over a RESTful API) and charged based on number of documents processed. However, they suffer from two key technical drawbacks. Firstly, document-by-document processing over HTTP is inefficient on large datasets and is also limited to within-document text processing algorithms. Secondly, the text processing algorithms are pre-packaged: it is not possible for researchers to extend the functionality (e.g. adapt such a service to recognise new kinds of entities). Additionally, these text processing SaaS sites come with daily rate limits, in terms of number of API calls or documents that can be processed. Consequently, using these services is not just limited in terms of text processing functionality offered, but also quickly becomes very expensive on large-scale datasets, especially web pages and tweets which tend to be short but numerous.

At the same time, NLP researchers have started developing distributed algorithms for data-intensive text mining over terabyte-scale datasets. Many approaches adopt the MapReduce model for distributed computation which can be deployed via Hadoop on clusters of affordable compute servers. For instance, Lin [15] demonstrates using Hadoop to calculate word co-occurrence matrices; later Lin and Dyer [16] present a number of distributed algorithms widely used in NLP tasks (Page-Rank, Expectation Maximisation, Hidden Markov Models) and discuss their MapReduce implementations.

¹ <http://www.opencalais.com>

² <http://extractiv.com>

³ <http://www.alchemyapi.com>

Meng [19] demonstrates using Hadoop with NLTK [17]; word counts, name similarity and $tf*idf$ are calculated and HMMs are trained. However, having to adapt NLP algorithms in this way is very time consuming. For example, van Gael and Bratieres from Cambridge University reported several person months of effort to learn Hadoop, Amazon AWS and to rewrite their part-of-speech tagging algorithm [9]. An added complication is that not all NLP algorithms are actually implementable in the MapReduce paradigm, e.g. those requiring shared global state.

Apart from needing new algorithms, large-scale text processing also requires access to sufficiently large clusters of servers. Utilising on-demand compute power is now possible through Infrastructure-as-a-Service (IaaS) providers [14]. Arguably the most successful commercial one is currently Amazon, who offers various kinds of virtual machine instances. However, making optimal use of virtual IaaS infrastructures for data-intensive NLP again comes with a significant overhead, e.g. having to learn the intricacies of Amazon’s APIs for the elastic compute cloud and simple storage services.

Platform-as-a-Service (PaaS) [5] are a type of cloud computing service which insulate developers from the low-level issues of utilising IaaS effectively, while providing facilities for efficient development, testing, and deployment of software over the Internet, following the SaaS model. In the context of traditional NLP research and development, and pre-dating cloud computing by several years, similar needs were addressed through NLP infrastructures, such as GATE [4] and UIMA [6]. These infrastructures have accelerated significantly the pace of NLP research, by providing a number of reusable algorithms (e.g. rule-based pattern matching engines, common machine learning algorithms), free tools for low-level NLP tasks (e.g. tokenisation, sentence identification), and in-built support for multiple input and output document formats (e.g. XML, PDF, DOC, RDF, JSON).

Some attempts have been made to extend UIMA functionality to utilise cloud computing; Zhou et al [23] propose a layered architecture built on UIMA and Hadoop with the intention of allowing UIMA to operate over very large datasets. Thus far, only a proof of concept on topic detection has been described. Ramakrishnan et al [20] explore the topic of scaling up processing for the SciKnowMine project. They discuss various possibilities, including cloud computing via the Behemoth project, which aims to make both UIMA and GATE functionality available on Hadoop. Luis and Matos [18] aim to make the benefits of cloud computing available outside of established NLP frameworks. They describe an approach to integrating various NLP tools and executing them in parallel. They comment that UIMA and GATE would benefit from adopting MapReduce. Laclavic et al [12] demonstrate using Ontea [11] with Hadoop.

This paper presents GATECloud.net – the adaptation of the GATE infrastructure to the cloud, following the PaaS paradigm. It enables researchers to run their NLP applications without the significant overheads of re-implementing their algorithms for MapReduce and understanding Amazon’s IaaS APIs.

GATECloud.net is novel and unique in that it is currently the only open-source cloud-based PaaS for large-scale text processing. Similar to the text processing SaaS discussed above, it offers a growing number of pre-packaged NLP services. However, due to being a specialised NLP PaaS, GATECloud.net also supports a bring-your-own-pipeline option, which can be built easily by reusing pre-existing NLP components and adding some new ones. Moreover, GATECloud.net is the only cloud-based NLP platform that supports *the complete NLP development lifecycle*. In addition to offering entity extraction services like OpenCalais, our NLP PaaS also supports data preparation

(e.g., HTML content extraction), manual corpus annotation, measurement of inter-annotator agreement, performance evaluation, data visualisation, indexing and search of full text, annotations, and ontological knowledge.

Crucially for researchers who only run large-scale text processing experiments infrequently, there are no recurring monthly costs: instead, GATECloud.net is pay-per-use, billed per hour (due to Amazon’s per-hour charging model). There is also no daily limit on the number of documents to process or on document size. Consequently, processing costs depend on the **total data size**, not on the number of documents. This is particularly advantageous for bulk processing of Twitter data and other such numerous, but small-sized texts. One downside of our per-hour billing approach is that it makes it harder for users to estimate likely usage costs upfront. We return to this issue in Section 5, where we discuss how users can overcome this problem.

3. Towards an NLP PaaS: Requirements and Methodology

NLP platforms, such as GATE and UIMA, have been hugely successful thanks to the clean separation between low-level tasks such as data storage, data visualisation, location and loading of components and execution of processes from the data structures and algorithms that actually process human language. They also reduce significantly the integration overheads by providing standard mechanisms for NLP components to be combined into complete pipelines and to communicate data about language, using open standards such as XML and RDF.

In a cloud computing context, developing an NLP PaaS requires the careful consideration of the following additional requirements:

- (i) *Straightforward deployment and sharing of NLP pipelines*: How can we achieve this transparently for the NLP developer, i.e. NLP applications developed on a desktop machine can run without any adaptation on the PaaS. In addition, developers need to be able to share easily their NLP pipelines as SaaS, with on-demand scalability and robustness ensured by the underlying NLP PaaS.
- (ii) *Efficient upload, storage, and sharing of large corpora*: An NLP PaaS needs to support a secure and efficient way for users to bulk upload, analyse, and download large text corpora, i.e. batch processing over large datasets. In addition, users need to be able to share their large text corpora between different NLP pipelines, running on the PaaS - both for services bundled within the NLP PaaS and services created by the developers themselves.
- (iii) *Algorithm-agnostic parallelisation*: how best to parallelise the execution of complex NLP pipelines, which could contain arbitrary algorithms, not all of which are implemented/suitable for MapReduce and Hadoop.
- (iv) *Load balancing*: determine the optimal number of virtual machines for running a given NLP application within the PaaS, given the size of the document collection to be processed and taking into account the considerable overhead of starting up new virtual machines on demand.
- (v) *Security and fault tolerance*: As with any web application, the NLP PaaS needs to ensure secure data exchange, processing, and storage, as well as be robust in face of hardware failures and processing errors.

In addition to these technical requirements, an NLP PaaS needs to offer comprehensive methodological support to underpin the NLP application development lifecycle:

1. Create an initial prototype of the NLP pipeline, testing on a small document collection, using an NLP application development environment, running on a standard desktop or a local server machine;
2. Crowd-source a gold-standard corpus for evaluation and/or training, using a web-based collaborative corpus annotation tool, deployed as a service on the PaaS;
3. Evaluate the performance of the automatic pipeline on the gold standard (either locally within the desktop development environment or through the manual annotation environment on the cloud). Return to step 1 for further development and evaluation cycles, as required.
4. Upload the large datasets and deploy the NLP pipeline on the PaaS;
5. Run the large-scale text processing experiment and download the results as XML, JSON, RDF, or schema.org formats. Optionally, an NLP PaaS could also offer scalable semantic indexing and search over the linguistic annotations and the document content.
6. Lastly, analyse any errors, and if required, iterate again over the required system development stages, either on a local machine or on the NLP PaaS.

Next, we present GATECloud.net – a fully implemented NLP PaaS and discuss how specifically we chose to address the technical and methodological requirements discussed above.

4. GATECloud.net: An Implemented NLP PaaS

(a) An Overview of the GATE NLP Infrastructure

The GATE framework [4] was chosen as the most suitable open-source NLP infrastructure to underpin the GATECloud.net PaaS, since it provides a unique combination of tools, addressing all the steps in the NLP application development methodology outlined above:

- A comprehensive and extensible NLP application development environment (GATE Developer), offering specialised user interfaces for visualisation and editing of linguistic annotations, parse trees, ontologies, and other NLP-specific resources (e.g. name lists, lexicons), as well as numerous tools for automating performance evaluation of language processing components.
- Numerous reusable text processing components for many natural languages, underpinned by a finite state transduction language (JAPE) for rapid prototyping and efficient implementation of shallow NLP analysis methods, as well as an extensible machine learning layer.

- A multi-paradigm repository (GATE Mímir) which can be used to index and search over text, linguistic annotations, semantic schemas (ontologies), and semantic meta-data. It supports queries that arbitrarily mix full-text, structural, linguistic and semantic constraints and scales up to terabytes of text through federated indexing.
- A web-based annotation environment (GATE Teamware) for collaborative creation of manually annotated corpora (needed for NLP algorithm training and for quantitative evaluation).

GATE is widely used for building text processing application in diverse domains, including bio-informatics [10], social media analysis [22] and humanities computing [21]. Recent projects have increasingly faced the problem of running GATE-based text processing on terabyte datasets, e.g. [2]. At the same time, the multi-tier service-oriented architecture of GATE Teamware, coupled with its centralised workflow engine, has made its deployment and administration too complex and error-prone for many researchers.

(b) *The GATECloud.net Architecture*

The high-level approach taken in GATECloud.net is to layer the GATE NLP infrastructure on top of Amazon’s cloud infrastructure and services, much in the same way in which GATE insulates researchers from having to deal with data formats, storage, and command-line processing, while also making it straightforward to create new and run text analysis algorithms.

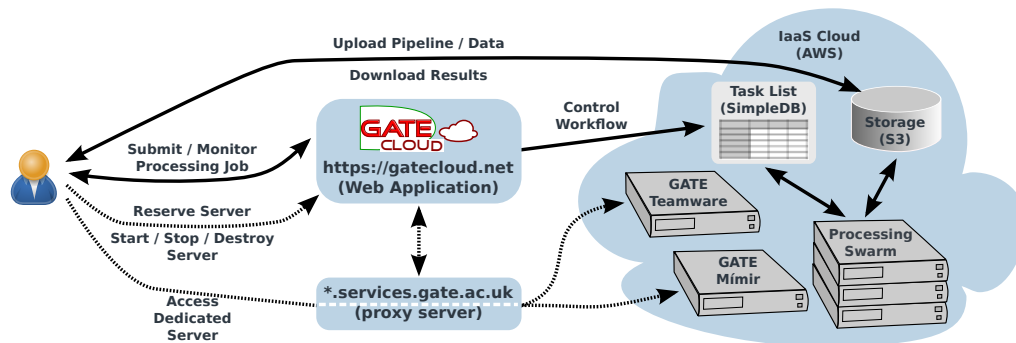


Figure 1. GATEcloud.net Architecture: dashed lines refer to supporting on-demand servers, while solid lines are for large-scale processing jobs

The GATECloud.net⁴ architecture comprises the following main elements:

GATECloud.net web application a web site implemented using the Grails⁵ web application framework. Its main functions are providing security and user authentication, web-based user interface for virtual servers management and annotation jobs management.

⁴ Available at <http://gatecloud.net>.

⁵ <http://grails.org>

Amazon Web Services as the cloud infrastructure provider. Elastic Compute Cloud (*EC2*) provides virtual server instances that are used for hosting on-demand servers and for text processing (annotation) jobs. *SimpleDB* is used for jobs and task management across distributed compute swarms. Finally, the Simple Storage Service (*S3*) stores software images for our on-demand servers, datasets, processing results, and processing reports in the case of annotation jobs.

services.gate.ac.uk a specially-configured Apache server that manages the *.services.gate.ac.uk sub-domain. We use that to provide persistent DNS names for the on-demand servers belonging to the system users. This insulates users from the implementation details of the Amazon cloud, where each server instance gets a new IP address and a new DNS name whenever it is started. Amazon EC2 offers a service named *Elastic IPs* where a set of static IP addresses can be allocated to the virtual servers as needed. However, this is limited to 5 addresses as standard, while we needed to be able to support far larger numbers of machine reservations.

<i>Making a server reservation</i> • The user requests a new reservation through the GATECloud.net web site, • a persistent data volume is created and associated with the user account, • a persistent server name is reserved for the user, • the user is notified in the web dashboard and through email.	<i>Destroying a server reservation</i> • The user requests the destruction of the reservation, • a check is made that the reservation is not currently associated with a running instance, • the previously created data volume is deleted, • the reserved server name removed from the database, • the user is notified.
<i>Starting a server</i> • The user requests the start-up of the server, • a new instance is started, • the previously created data volume is attached to the new instance, • the reserved server name is associated with the IP address of new instance, • the start-up process is confirmed to have completed successfully, • the user is notified.	<i>Stopping a server</i> • The user requests the stopping of the server, • the software services are stopped, • the persistent data partition is detached, • the server instance is terminated, • the user's account is updated with the associated costs, • the link between the reserved server name and the IP address of the terminated instance is removed, • the user is notified.

Table 1. On-demand server workflows

(c) On-demand Servers

As discussed above, the GATE family of NLP tools includes server-side platforms (GATE Teamware and GATE Mimir), both having complex architectures and sets of dependencies that make them difficult to install and maintain. In many NLP projects, these tools are only needed for a limited time only (e.g. Teamware is only required for manual annotation of training and evaluation corpora), which can make them uneconomical to install and maintain. A cloud-based SaaS deployment thus makes sense, as it reduces the costs in terms of admin effort and server hardware required.

Consequently, we created pre-installed Teamware and Mimir servers, running as Amazon EC2 instances, based on 64-bit Ubuntu Linux. Instead of virtual machine definitions (*Amazon Machine Images*, or *AMIs*) for server description, we chose to use ‘*recipes*’ instead. Such a *recipe* comprises scripts and configuration data describing how to transform a standard base *AMI* into a dedicated Teamware or Mimir server. These scripts include steps such as downloading and installing the software and the appropriate support packages (e.g. Java or MySQL), initialising the databases and configuration files, attaching the persistent data storage associated with the user owning the machine. These steps are executed automatically each time a new server instance is started.

One downside of our approach is that it slightly increases the start-up time, typically by less than five minutes. However, the advantages are manifold:

Reduced administration costs: since our dedicated servers are built on top of standard AMIs, the responsibility for maintaining the basic operating system (e.g. installing security updates, or upgrading software packages to newer versions, etc.) falls with the publisher of the AMI. We are currently using official Ubuntu Linux images, but any system that supports Java would work equally well.

Flexibility: upgrading to a new version of the base operating system (or even starting to use a completely different OS) is simply a matter of changing the base AMI identifier used in our recipe.

Improved security: each server instance is essentially re-installed from scratch on each start-up. This reduces the scope for possible persistent attack vectors.

In order to provision a server, a GATECloud.net user needs to first make a reservation for it. Once reserved, the server appears in the user’s dashboard, from where they can start it and stop it as required. While the server is running, the user incurs hourly charges. When the server is not required any more, the user can destroy the reservation, thus releasing all resources associated with it, such as the persistent data partition. The workflows used for these different operations are shown in Table 1.

(d) On-demand Large-Scale Text Processing

The other half of the GATECloud.net infrastructure is the support for processing of large document collections using cloud computing. As discussed in Section 3, there are five major technical requirements that need to be met, coupled with many methodological ones, arising from the specifics of text processing.

The implementation of *Annotation Jobs* on GATECloud.net addressed most of these technical and methodological requirements, leaving researchers free to concentrate on their experiments. From a researcher’s perspective, processing a document collection involves a few simple steps:

- upload the document collection (or point the system to the location where the data is available),
- upload a GATE-based processing pipeline to be used (or choose one provided as SaaS on GATECloud.net),
- press the ‘Start’ button.

While the job is running, a regularly updated execution log is made available in the user’s dashboard. Upon job completion, an email notification is also sent. Most of the implementation details are hidden away from the user, who interacts with the system through a web-based job editor, depicted in Figure 2.

Figure 2. Web-based Job Editor

Next we discuss how GATECloud.net PaaS meets the requirements from Section 3:

(i) Pipeline deployment and sharing: GATECloud.net can run any NLP pipeline developed on the researcher’s desktop, after it has been packaged by the GATE Developer environment [4], using the application export option “Export for GATECloud.net”. This is an automated process that builds a self-contained zip file, including all text processing modules and the linguistic data and ontologies required by them. In other words, there is no additional implementation effort required, in order to deploy a GATE-based NLP pipeline on the cloud-based NLP PaaS and execute it on a large-scale dataset.

(ii) Efficient Data Management: From an implementation perspective, job execution requires access to potentially large amounts of data, including the text processing application file(s), the document collection to be processed, the execution reports, the results files (if any are produced). We chose to store all this data in a separate S3 ‘location’ that is created especially for each job and cannot be accessed by other jobs. Once the job completes, users get a grace period (by default 10 days, but this can be changed) during which they can download the results. When this time has elapsed, the job’s S3 location is deleted entirely in order to save storage costs. Alternatively, the user can chose to provide their own S3 location for storing the job data (with the appropriate permissions for the GATECloud.net AWS user account), in which case no automatic deletion takes place.

(iii) Parallelisation and (iv) Load Balancing: Job execution is supported by multiple inter-connected, parallel workflows, which are described next.

Job Management: Each *Job* advances through a series of states during its life-cycle. The jobs management process is implemented by the GATECloud.net web application and deals with taking the appropriate steps to transition each active job from one state to the next. Each step usually involves generating new *Tasks*, and queueing them for execution. A *Task* is a piece of work that needs to be performed, and that is executed by a *Node* (a server instance in the cloud). Once all the *Tasks* belonging to a *Job* in a given state have completed, the *Job* can move to the next state.

Swarm Load Management: *Swarms* are sets of cloud instances (*Nodes*) that have identical configurations and that execute *Tasks* from the same task queue. Each *Swarm* has a target *load factor*, defined as $\frac{\text{active tasks} + \text{pending tasks}}{\text{active nodes}}$. A *load factor* of 1 indicates the intention of allocating a separate node for each running task, while a value of 2 means that the system would aim to have e.g 5 running nodes for a list of 10 tasks. The *Swarm* management process (also part of the the GATECloud.net web application) starts new *Node* instances associated with all of the registered *Swarms* as required in order to keep the actual *load factors* as close as possible to the target values.

Node Workflow: Each *Node* instance is configured during start-up with the tasks queue it should use. Once started, its workflow is a simple loop collecting the first pending *Task* from the queue, and executing it. When no more pending tasks have been available for a while, the node shuts itself down automatically. Each *Task* has an associated state and a time stamp. When a *Node* picks up a task for execution, it changes its state from *pending* to *active*. It also updates the time stamp at regular intervals while the *Task* is being executed. If a *Node* crashes for whatever reason, any *Task* it was working on will be left in an *active* state; however its time stamp will stop being updated. Active *Tasks* with an old time stamp are presumed to have failed and will be re-scheduled for execution three times. If a *Task* keeps failing to execute, then its state is changed to *failed* and it is not scheduled again, thus avoiding infinite loops.

The *Tasks* queue is implemented as an *Amazon SimpleDB* domain accessible by both the GATECloud.net web application and all processing *Nodes*. Concurrency control is implemented by associating a *version* attribute with each item in the domain. The *version* starts at 0 and is incremented with each write operation. Completing the picture are the atomic conditional update operations provided by *SimpleDB* which are used to make sure new values are only written if the *version* has not been independently changed since the old values were read.

(v) Security: Since the PaaS allows researchers to upload their own NLP pipelines, we have no control over the software included within: it could contain poorly written, insecure, or even malicious code. Consequently, security is a major concern and we are addressing this by judicious privilege isolation. The system user, which has access to cloud login credentials and system settings, is only executing code that has been produced by the platform authors and security audited. User data and user processes are executed by a separate OS-level user, which has very restricted permissions, allowing file access only within a given data directory. All the data and processes owned by this user are destroyed upon completion of each *Task*. This ensures that data and software belonging to different researchers are never present at the same time on the same *Node* machine instance.

While stored on the cloud, the user data is protected by the security procedures instituted by the cloud infrastructure provider. All the transfers between the cloud

storage, the processing nodes, and the user’s computer are done via an encrypted channel, using SSL.

5. Use Cases and Experiments

In order to quantify the performance gains from running text processing algorithms on large datasets on GATECloud.net, we carried out experiments on three document collections: 50 million tweets (very short texts), 20,000 news articles (medium-sized texts), and 100,000 patents (larger sizes, up to 6 MiB). With respect to data sizes, the statistics for the three datasets are as follows:

Twitter We randomly selected 50 million tweets from a 1TB dataset. The Twitter feed was first simplified by only preserving the tweet text, tweet ID and the ID of the author, while discarding all other metadata fields. After conversion to XML, the size of the resulting corpus was 6 GiB. The 50 million tweets were chosen to be a usefully sized dataset, 10 times larger than the one used by Abel *et al* [1] and 10 times smaller than the one used by Laniado and Mika [13]. Additional experiments, not detailed here due to space constraints, showed that processing time scales linearly with the number of tweets, on all hardware configurations (see below). Consequently, the performance and cost figures reported here can be used to easily estimate the corresponding values for smaller or larger tweet datasets.

News The news corpus comprised 20,000 HTML pages (1.31 GiB), collected from the web sites of news broadcasters *BBC* and *CNN*, and UK newspapers *The Independent* and *The Guardian*. The shortest document was just 9 characters (from the CNN web site), whereas the longest was 230 KiB. Most news articles are well clustered around the middle, with an average size of 68.7 KiB.

Patents The corpus contained 100,000 patent documents, with an overall size of 5.47 GiB. The mean document size is around 58 KiB, with the smallest document being just under 3 KiB (containing only the abstract) and the largest at 5.94 MiB. The majority of the patent documents are clustered around the middle, with lower quartile: 22 KiB, median of 43 KiB, and the upper quartile: 70 KiB.

The news and Twitter datasets were annotated for named entities with the standard *ANNIE* entity annotation pipeline [4], deployed as SaaS within GATECloud.net. For the patents dataset we reused a pre-existing text processing pipeline [2] which recognises patent-specific types, including references to other patents, scientific publications, measurement expressions, patent sections, claims, examples, references to figures and tables.

In order to carry out a cost-benefit evaluation of GATECloud.net, we ran three sets of experiments, where each dataset was processed on the following hardware configuration:

Desktop *Lenovo ThinkCentre M58p* desktop computer, Intel Xeon E5502 1.86 GHz CPU (2 cores), 4 GiB RAM, 320 GB HDD Serial ATA II, cost $\approx$ £1300. This is currently a standard desktop configuration for researchers in our laboratory.

Server *HP ProLiant DL385 G7* server, AMD Opteron 2.3 GHz CPU (12 cores), 32 GiB RAM, 2 TB disk space, cost $\approx$ £4800. For these experiments, we had access to 6 of the processing cores only.

Cloud *GATECloud.net* swarm, using 10 *Extra Large* instances on the Amazon Cloud. Each node has 15 GiB of RAM, 8 EC2 Compute Units (4 virtual cores with 2 EC2 Compute Units each), 1690 GB of local instance storage, 64-bit platform. One EC2 Compute Unit (ECU) provides the equivalent CPU capacity of a 1.0-1.2 GHz 2007 Opteron or 2007 Xeon processor. For running costs, see Table 2.

The only exception is the patents dataset which was too large to be processed on the desktop computer in a reasonable amount of time.

		Time (hh:mm:ss)			Speed (KiB/s)	Cost (GBP)	
		CPU Time	Computer Time	Clock Time		£/GiB	Total
Experiment 1: Patents 100,000 patent documents	Desktop	N/A	N/A	N/A			
	Server	91:42:00	18:39:54	18:39:54	85.33		
	Cloud	162:38:00	16:56:37	02:03:32	773.6	£3.08	£16.83
Experiment 2: News 20,000 documents	Desktop	05:20:19	05:20:19	05:20:19	71.52		
	Server	04:43:00	03:08:00	03:08:00	121.86		
	Cloud	07:47:00	01:21:20	00:35:31	645.04	£1.51	£1.98
Experiment 3: Tweets 50,000,000 tweets	Desktop	32:28:46	32:28:46	32:28:46	52.80		
	Server	22:16:15	03:19:12	03:19:12	516.53		
	Cloud	40:08:00	07:00:14	01:25:46	1199.69	£1.35	£7.92

Table 2. Experiment Results

Table 2 reports the results of these three sets of experiments along several dimensions. *CPU Time* measures the total amount of time spent, by any thread, processing documents, while *Computer Time* is the time taken by a given machine. On machines with multiple cores, Computer Time will be lower than the corresponding CPU time, as multiple parallel execution threads are used. *Clock time* represents the time interval between starting the process and its completion, as measured by an external clock. This differs significantly from the other measurements only in the case of the cloud-based experiments, where processing was distributed across several machines. On the desktop machine, even though it has two cores, only one of them was used due to running the experiments as standard batch processes, using a single execution thread. Consequently, the desktop running times are always the same for CPU time, computer time, and clock time. The *Clock time* column shows the overall time saving due to using more powerful hardware configurations: from a desktop computer to a server, then to the cloud.

In terms of cost, the price per GiB for named entity recognition is very low (between £1.35 and £1.50). The patents application is significantly more complex and therefore, its cost per GiB processed is approximately double. Similarly, its performance figures should not be compared directly to those obtained on news and tweets. The more general point is that the cost per GiB of text processed on GATECloud.net varies widely depending on the complexity of the underlying text processing algorithms. In addition, specifically on the patents dataset, significant amount of time was spent on processing the few larger documents of over 1 MiB in size. Time complexity for co-reference resolution is higher than linear, so it tends to be much more time consuming on such large documents due to the large number of candidate entities that need to be checked. The time-complexity of most other algorithms tends to be near linear with respect to document size.

It is, of course, required that users are able to estimate the cost of processing a given document collection. Given that each document is processed independently from the

rest, the time taken to process a collection can be approximated by multiplying the total number of documents and the average execution time for a document⁶. The average execution time can be estimated by processing a small but statistically representative collection sample. Care should be taken when selecting the sample as, depending on what the actual processing consists of, document length may not be a good indicator of time complexity.

Compared to cloud infrastructure providers, GATECloud.net is simplifying cost estimations by hiding details such as the cost of data storage and data transfers, behind a simple cost model based exclusively on CPU-hours.

When looking at CPU time used, this always tends to be higher on the cloud than the CPU times for the desktop and server. This is due to the overhead in using distributed computing, specifically mainly due to the need to split the large datasets into batches that run on separate nodes, as well as efficiency lost to virtualisation. In addition, the actual hardware specifications of the Amazon *Extra Large* virtual machines are not as good as those of the server, but are quite comparable to our desktop configuration.

As can be seen from the time statistics, the major benefit of using GATECloud.net comes from the significant reduction in *Clock Time* taken for each of the experiments. The speed-up compared to the Desktop configuration is between 10- and 20-fold. For example, the processing time for the news dataset is reduced down from over 5 hours to 35 minutes, which can help significantly not only for processing large-scale datasets, but also during the development of the algorithms by reducing the time taken by the develop-evaluate cycles. The time reduction is even greater on the tweets dataset, where time goes down from 32 hrs to 1 hr and 25 minutes.

With respect to the gains made by using GATECloud.net instead of a local powerful server, there are also significant benefits (5-fold speed-up on the news set and 10-fold on the patents data). The benefits are less pronounced on tweets, due to their large number and small size. In general, GATECloud.net has been optimised for processing medium- to large-sized documents, where the benefits are most pronounced. In future work, we will be working towards improving the infrastructure's performance on large collections of smaller documents.

An independent GATECloud.net benchmarking experiment was carried out by a team of researchers from the UK Food and Environment Research Agency⁷. They started by building a specialised text annotation pipeline, using GATE Developer. This included over 20 different rule-based components, as well as some of the low-level linguistic processing offered by GATE's standard tools. The document collection consisted of 261,260 documents ranging from 1 KB to 2.5 MB of text. Of all the documents, 14 (i.e. 0.005%) failed to complete successfully, due to various text processing exceptions. Their *CPU time* was just over 13 hours (786 minutes), whereas the *Clock time* value was 1 hr and 20 minutes, again a 10-fold speed-up.

6. Conclusions and Future Work

This paper motivated the need for a specialised NLP PaaS, identified a set of requirements, and presented our implementation of such an NLP PaaS –

⁶ The actual formula is more complex, but this is the dominant term.

⁷ <http://www.fera.defra.gov.uk/>

GATECloud.net. This is aimed at helping researchers to carry out data-intensive text processing experiments on the cloud.

The platform has been made available to the public as a beta service. During its first 6 months of operation (between mid June and mid December 2011), there were 114 registered users. The number of processed documents was 4.7 million⁸, part of 302 annotation job runs. The accumulated server time used for both annotation jobs and dedicated servers was 430 hours. This level of usage indicates a need for such tools and a clear interest from researchers and the wider community.

Currently we are working on adding a programming API to GATECloud.net, so that data upload, processing, and download can all be done automatically, outside of the web interface. This will allow tighter integration with legacy work-flows and higher levels of automation. Other future work will focus on going beyond batch-oriented processing of documents, towards handling streaming data. Integration of MapReduce-based NLP algorithms is also being considered.

Acknowledgment

This work has been partially supported by a EPSRC/JISC grant (EP/I034092/1), funding pilot projects in Cloud Computing. Kalina Bontcheva is supported by a Career Acceleration Fellowship from the Engineering and Physical Sciences Research Council (EP/I004327/1).

References

- [1] F. Abel, Q. Gao, G.-J. Houben, and K. Tao. Semantic enrichment of twitter posts for user profile construction on the social web. In *ESWC (2)*, pages 375–389, 2011.
- [2] M. Agatonovic, N. Aswani, K. Bontcheva, H. Cunningham, T. Heitz, Y. Li, I. Roberts, and V. Tablan. Large-scale, parallel automatic patent annotation. In *Proc. of 1st International CIKM Workshop on Patent Information Retrieval - PaIR'08*, Napa Valley, California, USA, October 30 2008.
- [3] R. Barga, D. Gannon, and D. Reed. The client and the cloud: Democratizing research computing. *IEEE Internet Computing*, 15(1):72–75, 2011.
- [4] H. Cunningham, D. Maynard, K. Bontcheva, V. Tablan, N. Aswani, I. Roberts, G. Gorrell, A. Funk, A. Roberts, D. Damjanovic, T. Heitz, M.A. Greenwood, H. Saggion, J. Petrak, Y. Li, and W. Peters. *Text Processing with GATE (Version 6)*. <http://gate.ac.uk/books.html>, 2011.
- [5] M. D. Dikaiakos, D. Katsaros, P. Mehra, G. Pallis, and A. Vakali. Cloud computing: Distributed internet computing for it and scientific research. *IEEE Internet Computing*, 13(5):10–13, 2009.
- [6] D. Ferrucci and A. Lally. UIMA: An Architectural Approach to Unstructured Information Processing in the Corporate Research Environment. *Natural Language Engineering*, 2004.

⁸ For the tweets experiments, we created input documents from groups of 1,000 tweets. Hence the 50,000,000 tweets only count as 50,000 documents.

- [7] I. Foster. Globus online: Accelerating and democratizing science through cloud-based services. *IEEE Internet Computing*, 15(3):70–73, 2011.
- [8] A. Halevy, P. Norvig, and F. Pereira. The unreasonable effectiveness of data. *IEEE Intelligent Systems*, 24(2):8–12, 2009.
- [9] M. Hammond, R. Hawtin, L. Gillam, and C. Oppenheim. Cloud computing for research – Final Report. Technical report, JISC, 2010. <http://www.jisc.ac.uk/whatwedo/programmes/researchinfrastructure/usingcloudcomp.aspx>.
- [10] M. Johansson, Y. Li, J. Wakefield, M. Greenwood, T. Heitz, I. Roberts, H. Cunningham, P. Brennan, A. Roberts, and J. McKay. Using prior information attained from the literature to improve ranking in genome-wide association studies. In *The American Society of Human Genetics 59th Annual Meeting*, Honolulu, Hawaii, USA, October 2009.
- [11] M. Laclavik, M. Seleng, and L. Hluchy. Towards Large Scale Semantic Annotation Built on MapReduce Architecture. In Marian Bubak, Geert van Albada, Jack Dongarra, and Peter Sloot, editors, *Computational Science – ICCS 2008*, volume 5103 of *Lecture Notes in Computer Science*, pages 331–338. Springer Berlin / Heidelberg, 2008.
- [12] Michal Laclavik, Martin Seleng, Emil Gatia, Zoltan Balogh, and Ladislav Hluchy. Ontology based Text Annotation – OnTeA. In *Proceedings of the 2007 conference on Information Modelling and Knowledge Bases XVIII*, pages 311–315, Amsterdam, The Netherlands, The Netherlands, February 2007. IOS Press.
- [13] David Laniado and Peter Mika. Making sense of twitter. In Peter Patel-Schneider, Yue Pan, Pascal Hitzler, Peter Mika, Lei Zhang, Jeff Pan, Ian Horrocks, and Birte Glimm, editors, *The Semantic Web – ISWC 2010*, volume 6496 of *Lecture Notes in Computer Science*, pages 470–485. Springer Berlin / Heidelberg, 2010.
- [14] A. Li. Comparing public cloud providers. *IEEE Internet Computing*, 15(2):50–53, 2011.
- [15] J. Lin. Scalable language processing algorithms for the masses: a case study in computing word co-occurrence matrices with mapreduce. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing, EMNLP '08*, pages 419–428, Stroudsburg, PA, USA, 2008. Association for Computational Linguistics.
- [16] Jimmy Lin and Chris Dyer. *Data-Intensive Text Processing with MapReduce*, volume 3. Morgan& Claypool, San Francisco, USA, 2010.
- [17] E. Loper and S. Bird. Nltk: The natural language toolkit. *CoRR*, cs.CL/0205028, 2002.
- [18] T. Luís and D. M. de Matos. High-performance high-volume layered corpora annotation. In *Proceedings of the Third Linguistic Annotation Workshop, ACL-IJCNLP '09*, pages 99–107, Stroudsburg, PA, USA, 2009. Association for Computational Linguistics.

- [19] X. Meng. Tutorial on MapReduce Framework in NLTK. Technical report. http://nltk.googlecode.com/svn/trunk/nltk_contrib/nltk_contrib/hadoop/doc/.
- [20] C. Ramakrishnan, W. A. Baumgartner, J. A. Blake, G. A. P. C. Burns, K. Bretonnel Cohen, H. Drabkin, J. Eppig, E. Hovy, C. N. Hsu, L. E. Hunter, T. Ingulfsen, H. R. Onda, S. Pokkunuri, E. Riloff, C. Roeder, and K. Verspoor. Building the scientific knowledge mine (SciKnowMine): a community-driven framework for text mining tools in direct service to biocuration. In René Witte, Hamish Cunningham, Jon Patrick, Elena Beisswanger, Ekaterina Buyko, Udo Hahn, Karin Verspoor, and Anni R. Coden, editors, *New Challenges for NLP Frameworks (NLPFrameworks 2010)*, LREC 2010, pages 9–14, Valletta, Malta, May 2010. ELRA.
- [21] T. Risse, S. Dietze, D. Maynard, and N. Tahmasebi. Using Events for Content Appraisal and Selection in Web Archives. In *Proceedings of DeRiVE 2011: Workshop in conjunction with the 10th International Semantic Web Conference 2011*, volume 779, pages 98–107, Bonn, Germany, October 2011. Marieke van Erp, Willem Robert van Hage, Laura Hollink, Anthony Jameson, Raphaël Troncy.
- [22] H. Saggion and A. Funk. Extracting opinions and facts for business intelligence. *RNTI Journal*, E(17):119–146, November 2009.
- [23] Bin Zhou, Yan Jia, Chunyang Liu, and Xu Zhang. A distributed text mining system for online web textual data analysis. In *Cyber-Enabled Distributed Computing and Knowledge Discovery (CyberC), 2010 International Conference on*, pages 1–4, Los Alamitos, CA, USA, October 2010. IEEE Computer Society.