

BIG DATA

Gruter's
빅데이터 플랫폼 아키텍처 및 솔루션 소개

2012.10

김형준



이 저작물은 크리에이티브 커먼즈 코리아 저작자표시-비영리-변경금지 2.0
대한민국 라이선스에 따라 이용하실 수 있습니다.

오늘의 주제

그루터는 뭘로 먹고 사나?

그루터의 내부 시스템은 어떻게 되어 있나?

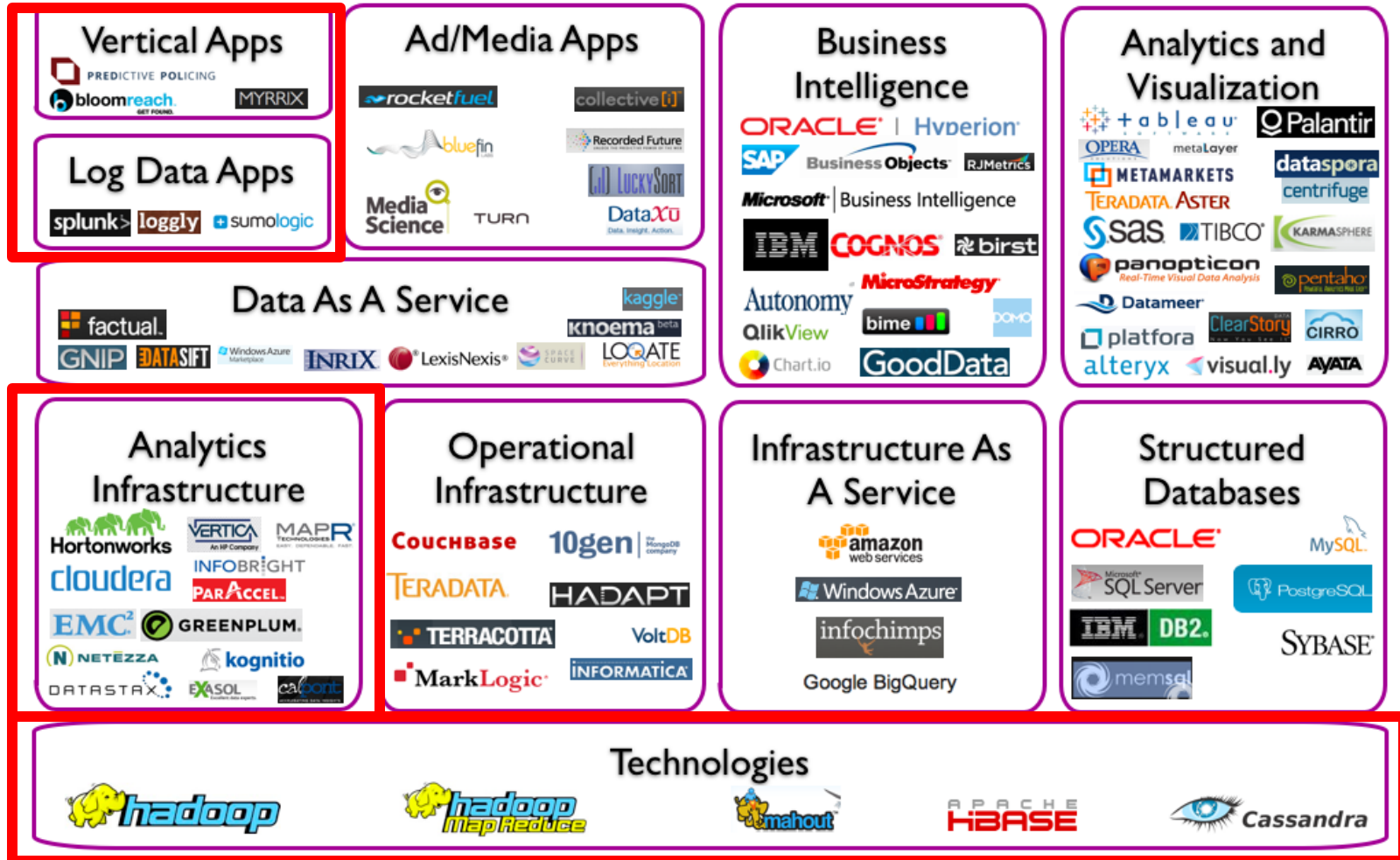
외부로는 어떤 서비스/솔루션을 제공하고 있나?

BIG DATA Market Segments

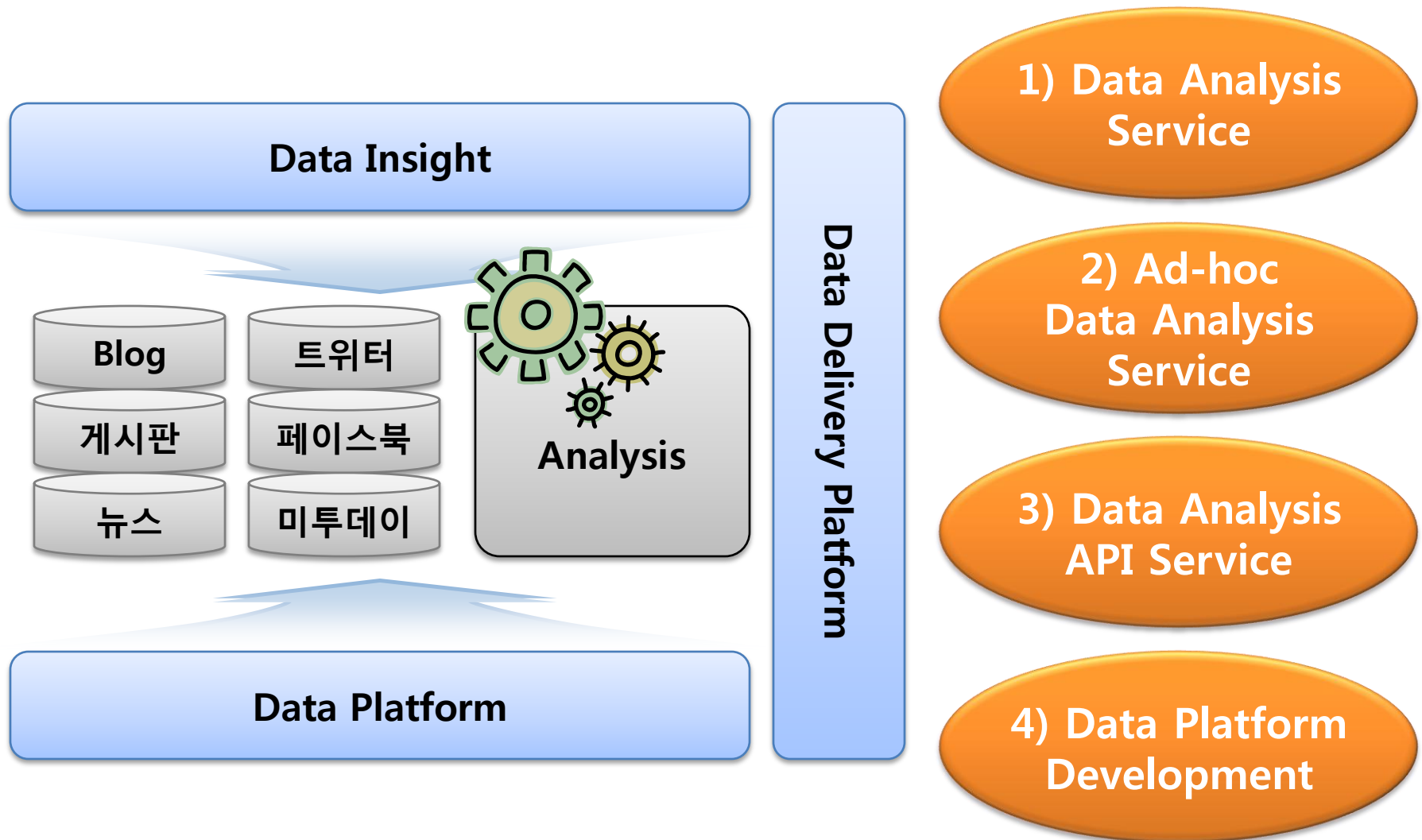
Hardware	Big Data Distributions	Data Management Components	Analytics and Visualizations	Services
- Storage - Servers - Networking Vendors include Dell, HP, IBM, Cisco.	- Community Hadoop distributions - Enterprise Hadoop distributions - Non-Hadoop Big Data frameworks Vendors include Cloudera, IBM, MapR, LexisNexis, Microsoft.	- NoSQL databases - Data integration - Data quality and governance Vendors include DataStax, IBM, Informatica, Syncsort.	- Analytic development platforms - Advanced analytics applications - Data visualization tools - Business intelligence applications Vendors include Karmasphere, Tresata, Datameer, SAS Institute, Tableau, Revolution Analytics.	- Consulting - Training - Software maintenance - Hardware maintenance - Hosting/cloud Vendors include Think Big Analytics, Amazon Web Services, Accenture, as well as services associated with enterprise distributions (e.g. Cloudera).
Next Generation Data Warehouse - MPP, columnar data warehouse appliances. - In-memory analytics engines Vendors include EMC Greenplum, HP Vertica, Teradata Aster Data, IBM Netezza, SAP, Microsoft, Kognitio.				

<http://wikibon.org/blog/navigating-the-big-data-vendor-landscape/>

Big Data Landscape



그루터의 비즈니스 영역



DATA PLATFORM DEVELOPMENT

BigData 프로젝트의 어려움

대상 데이터는?

데이터 관리 체계 부족
데이터 중심이 아닌 프로세스 중심의 운영
조직 내 데이터 권력이 존재

무엇을 분석하지?

도메인 전문가의 부족
스스로도 무엇을 분석해야 할 지 모름

시스템 구축에 대한
ROI는?

대상 데이터, 분석 대상이 모호
데이터 처리 시스템은 고가의 시스템들로 구성



2년째

계획 수립,
스터디만 계속

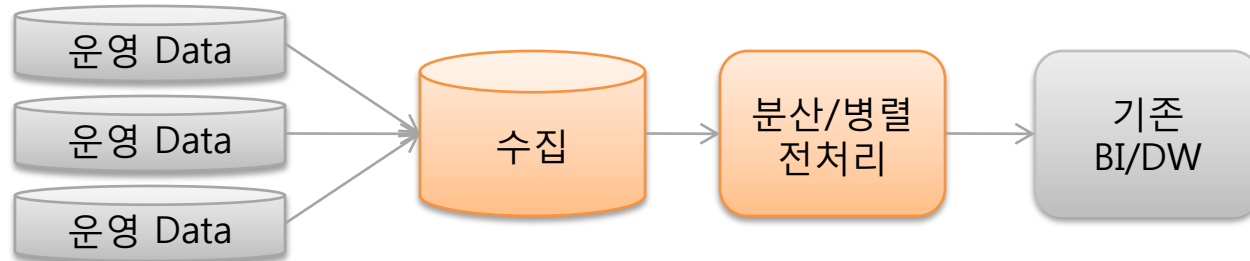
BigData 프로젝트 유형

기존 BI/DW 시스템 대체

- 기존 시스템의 분석 기능 대부분 제공해야 함
- 시스템 자체 다양한 분석 모듈 및 다양한 리포팅 지원
- 새로운 기술보다는 기존의 엔진을 업그레이드 하는 전략이 바람직

기존 BI/DW 시스템 보완

- 기존 시스템은 유지하면서 데이터 전처리 과정을 새로 구축
- 기존 시스템은 분석에 집중, 데이터 처리는 새로운 기술 활용



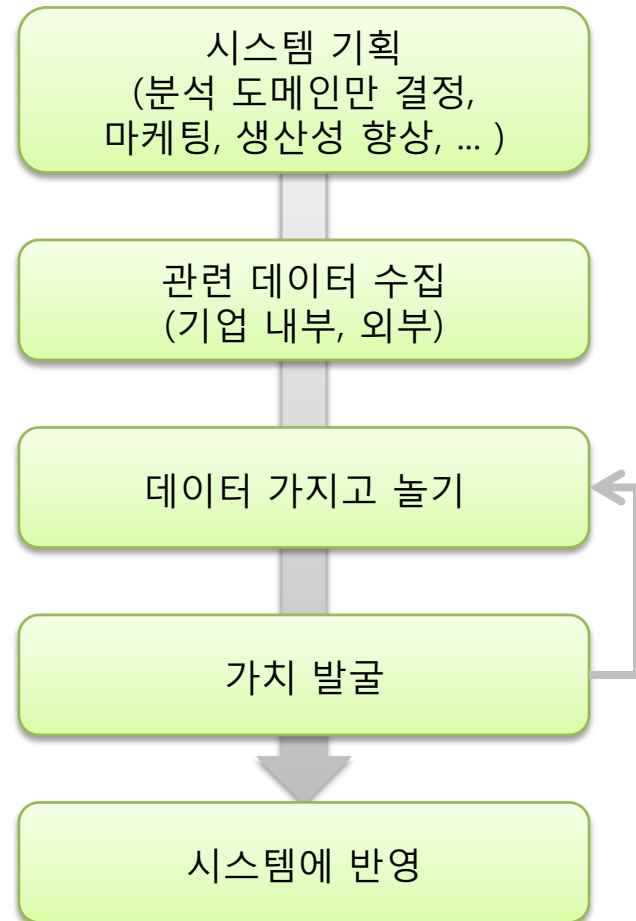
신규 데이터 처리 시스템 구축

- 오픈 소스 중심의 아키텍처로 구축
- 필요에 따라 상용 BI/DW 엔진 도입 및 연동

BigData 프로젝트 접근 방법



3 ~ 6개월 이상 소요



지속적인 활동

BigData 해결 방법은?

기존 데이터 처리 방식은 기존대로 활용

**새로운 데이터 처리에 대한 요구사항은
새로운 방식으로**

→ 데이터 민주주의 실현

누구나 쉽고 빠르게 데이터에 접근, 분석할 수 있는
체계(프로세스, 룰, 플랫폼 등) 구축

→ 적절한 팀 구성

도메인 전문가, 도메인 + 시스템 병행 전문가, 시스템 전문가

→ 시작은 작지만 지속, 확장 가능한 형태

ROI 및 조직 내 성과에 대한 고려

분석 항목, 데이터 증가 시에도 자연스럽게 개선 가능

빅데이터 플랫폼의 필요성

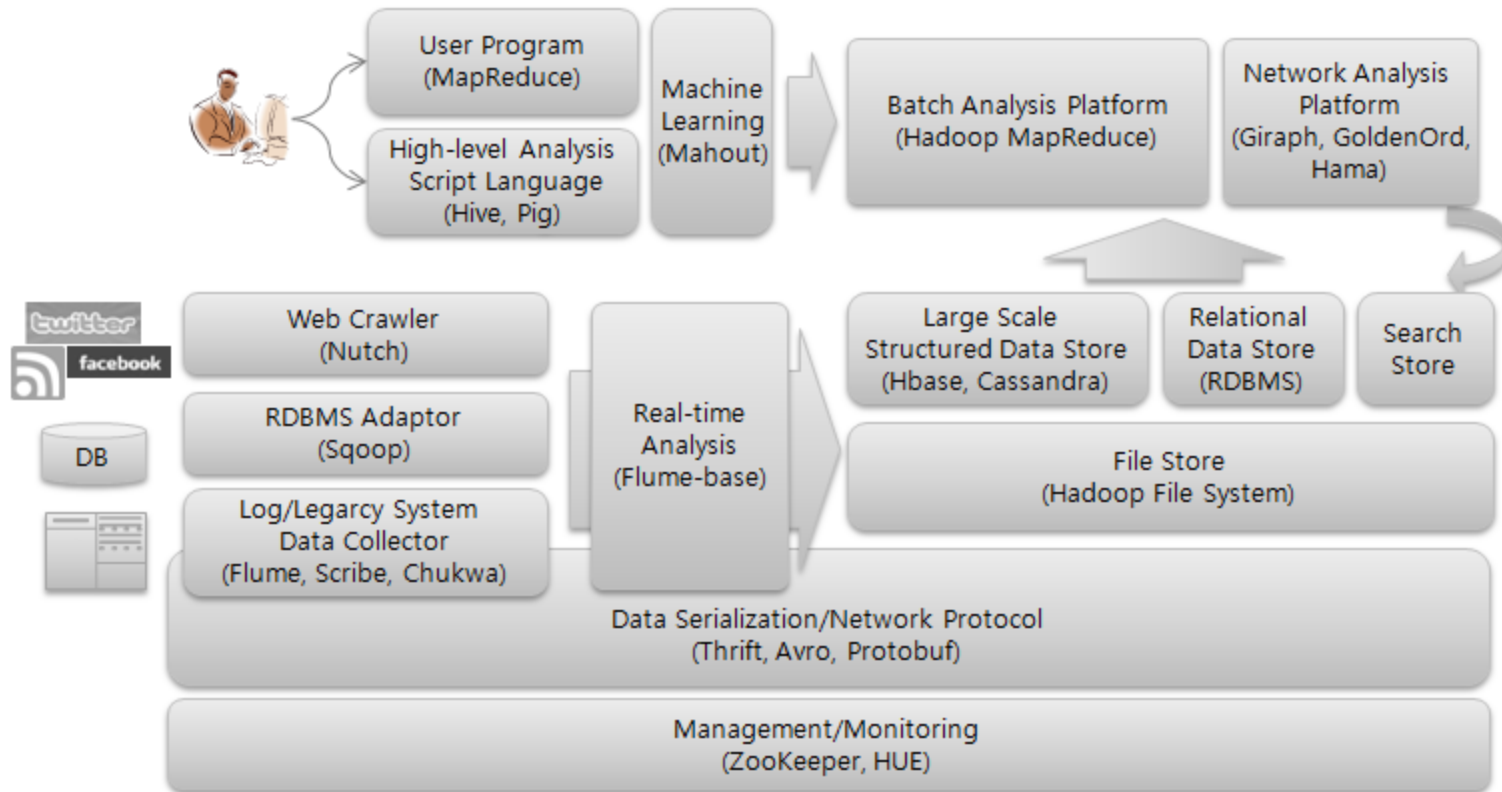
- 데이터의 다양성
 - 다양한 분석을 위해서는 기업의 내부 시스템뿐만 아니라 블로그, SNS 등과 같은 기업 외부의 공개된 다양한 종류의 데이터가 필요
 - 분석의 요구사항이 변함에 따라 기업 내부의 데이터에서도 필요한 데이터도 계속 변화
- 데이터 접근성
 - 무엇을 분석해야 할 지 모르는 경우가 많기 때문에 데이터 접근이 쉬워야 하며 수시로 분석 로직을 수행할 수 있어야 함.
 - 시스템 전문가가 아닌 데이터의 특징을 잘 알고 있는 데이터 전문가에 의한 분석 필요
 - 데이터 전문가가 직접 데이터에 접근해서 쉽고, 빠르게 분석할 수 있는 기반 제공 필요.
- 비용, 기술
 - 분석 요구사항 또는 데이터가 변할 때마다 시스템을 구축할 경우
 - 비용도 많이 들고
 - 시스템 구축 기간도 오래 걸리며
 - 기술도 분산되어 내재화 하기 어려움
 - 빅데이터 플랫폼을 기반으로 하여
 - 플랫폼 내에 다양한 데이터를 저장
 - 이를 활용하는 기능을 플랫폼 위에 구축함
 - 비용을 줄일 수 있으며
 - 기술력도 집중되는 효과

빅데이터 플랫폼 요구사항

- 데이터의 전체 라이프 사이클(수집, 저장, 분석, 폐기)을 관리하는 시스템
- 데이터 유형 변화에도 시스템의 변경 없이 적용, 운영 가능
- 다양한 분석 알고리즘 또는 분석 플랫폼이 적용 가능
 - Map/Reduce, MPI, Graph 등
- 비즈니스 요구사항에 부합되는 적절한 분석 Latency 지원
 - 실시간, 준-실시간, 배치
- 데이터의 용량 증가에도 즉시 대응 가능

Hadoop 기반 빅데이터 플랫폼

- 기본 아키텍처



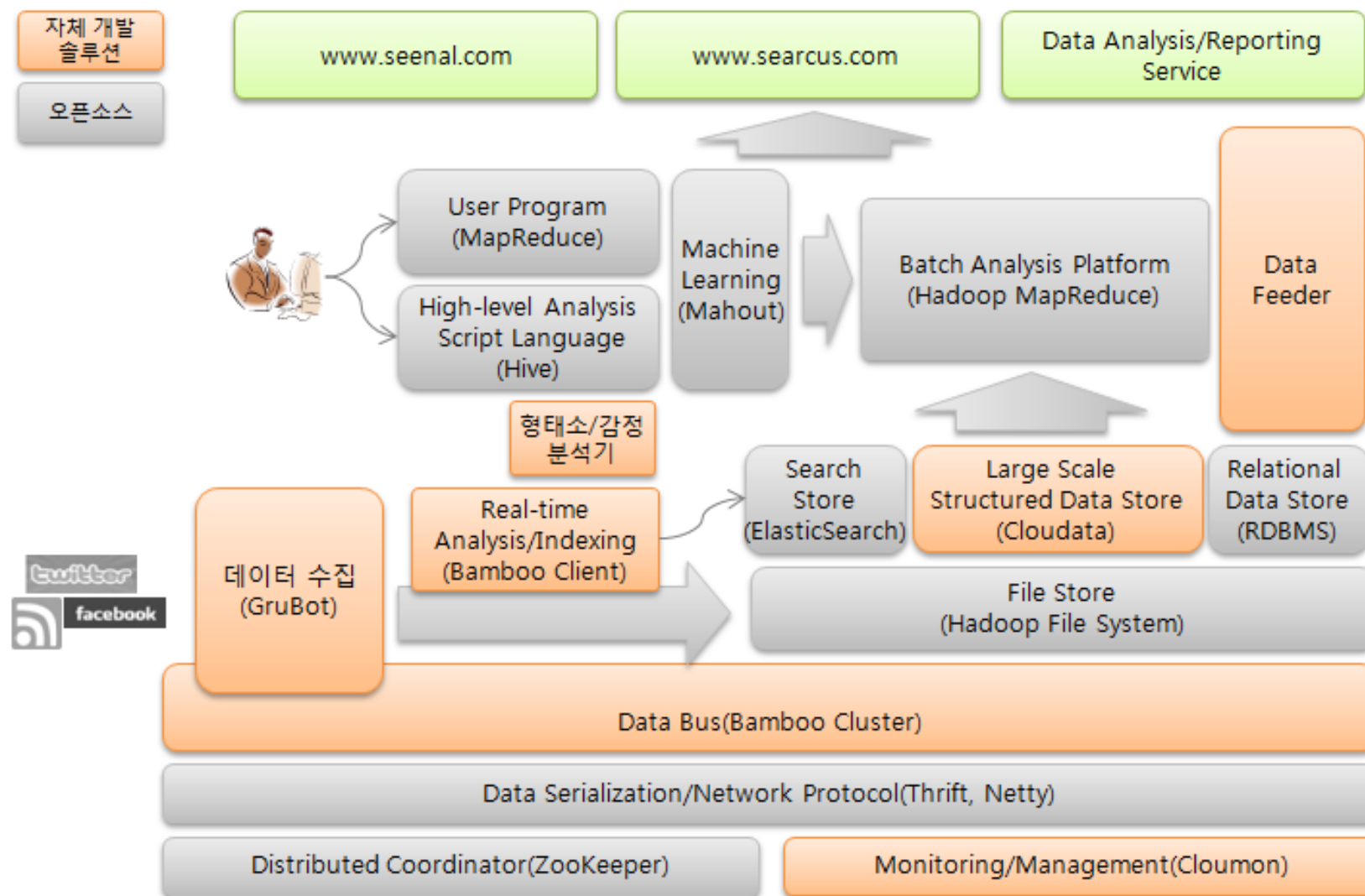
빅데이터 플랫폼 구성 시 주의 사항

- 요구사항에 적합한 솔루션 선정
 - 아키텍처를 구성하는 각 컴포넌트에 사용할 수 있는 솔루션은 하나가 아닌 경우가 대부분
 - 이 경우 구축하고자 하는 시스템의 요구사항에 가장 부합되는 솔루션이 무엇인지 검토하여 최적의 솔루션을 선정해야 함
- 기술적 이해
 - 빅데이터를 처리하는 솔루션의 대부분은 하나의 서버에서 운영되지 않고 대부분은 여러 대의 서버가 하나의 클러스터로 구성되는 분산 아키텍처
 - 일부 서버 장애 상황에 대한 처리, 메타 데이터의 관리, 서버 추가 시 메커니즘, 데이터 정합성 관리 등에 대한 기술적 이해가 필수
- 컴포넌트간 연결
 - 하둡 에코 시스템 기반 빅데이터 플랫폼은 다수의 독립적인 솔루션의 조합으로 구성
 - 다양한 솔루션을 유연하면서도 견고하게 구축하는 것이 중요
- 통합된 관리 및 모니터링 도구
 - 하둡 에코 시스템의 각 단위 솔루션은 좋은 관리 도구를 제공하지 않음
 - 각각의 솔루션을 조합하여 하나의 플랫폼을 구축할 경우 통합적 관리 도구 확보가 필수

qoobah: 그루터 빅데이터 플랫폼

- Hadoop eco system 기반 최적화된 빅데이터 플랫폼
 - 각 오픈 소스 솔루션에 대한 이해를 기반으로 한 아키텍팅
 - 각 오픈 소스의 장점, 취약점 등을 이해
 - 소스 코드 분석 수준의 기능 검증 및 분석
 - 오픈 소스 자체 미 제공 기능 및 취약 기능 해결 솔루션 제공
 - Gruter 자체 개발 모듈 제공
- 오픈 소스 기반 빅데이터 플랫폼 토털 서비스 제공
 - 오픈 소스 컨설팅
 - 기본 제공 오픈 소스에 대한 교육 및 운영 컨설팅
 - Hadoop, Hive, HBase, Flume, ZooKeeper 등
 - 오픈 소스에서 제공하지 않는 기능에 대해 그루터 개발 솔루션 제공
 - 실시간 분석 등 미 지원 기능 제공
 - 수집 시 분석 모듈과 연계 기능 제공 등
 - 기업 내 적용에 필요한 다양한 기능 제공
 - 데이터 수집 -> 분석의 전반적인 라이프사이클 관리 등
 - 오픈 소스의 취약한 부분을 해결하는 그루터 개발 솔루션 제공
 - 취약점 해결 방안 제시 및 자체 개발 솔루션 탑재하여 제공
 - 고객의 요구사항에 적합한 아키텍처 선정
 - 고객의 요구사항에 맞는 아키텍처 및 오픈 소스 선정
 - 고객 요구사항에 맞추어 오픈 소스 모듈 커스터마이징
 - 오픈 소스의 핵심 모듈은 수정하지 않으며 plugin 모듈 수정 및 추가 개발을 통해 고객 요구사항 만족

qoobah: 기본 아키텍처



qoobah 특징(1)

- 통합 데이터 관리 체계 제공
 - 빅데이터 시스템을 위한 데이터의 수집, 저장, 분석에 이르는 데이터 관리의 모든 단계를 지원하는 통합 플랫폼이다.
- 확장성
 - 대부분의 컴포넌트는 다수의 서버로 구성되어 있으며 데이터의 증가, 분석 업무의 증가, 분석 시간 단축 등과 같은 요구사항에 쉽게 확장할 수 있는 구조로 되어 있다.
 - 확장을 위해서 프로그램 또는 시스템 구성을 변경할 필요가 없으며 서버 추가만으로 쉽게 확장할 수 있는 구조를 가지고 있다.
- 안정성
 - 플랫폼을 구성하는 대부분의 컴포넌트는 이중화 이상으로 구성되어 있어 특정 컴포넌트의 일부 서버 장애 시에도 데이터의 수집, 분석, 서비스 제공은 정상 운영되도록 구성되어 있다.
 - 시스템을 지탱하는 중요 메타 데이터도 안정적으로 이중화 이상으로 저장한다.
- 유연성
 - 각 기능을 수행하는 오픈 소스 솔루션은 사용자의 요구사항의 맞추어 적합한 다른 대안 오픈 소스로 대체 가능하며 다른 솔루션으로 대체 시 유연하게 대응하도록 구성되어 있다.
- 빠른 성능
 - 대부분의 컴포넌트는 분산 아키텍처를 이용하여 수십 대의 서버에 분산 배치되어 있으며 이를 유기적으로 결합하여 분석 작업을 수행함으로써 빠른 분석 성능을 가지고 있다.

qoobah 특징(2)

- 데이터 접근 편의성
 - 저장된 데이터에 대한 분석은 개발자나 시스템 전문가가 아닌 데이터 전문가에 의해 수행되어야 하며 이를 위해서 데이터에 접근하기 쉬워야 하고 분석 작업 수행도 쉬워야 한다.
 - Qoobah는 원격에 있는 서버의 데이터의 수집부터 스토리지에 어떤 저장소로 저장되어야 하는지를 웹 기반 관리 UI를 이용하여 쉽게 설정할 수 있으며, 플랫폼에 저장된 데이터 역시 웹 기반 관리 UI를 이용하여 쉽게 조회하고 간단한 질의 언어를 이용하여 분산 환경에서 빠르고 쉽게 분석할 수 있다.
- 저렴한 구축 비용
 - 플랫폼의 주요 구성 요소는 대부분 하둡 에코 시스템이나 오픈 소스 위주로 구축되어 있고 사용하는 하드웨어 장비도 x86기반의 리눅스 장비를 이용하여 시스템 구축 및 확장 비용이 저렴하다.
- 독자적인 기술력 확보
 - Qoobah는 대부분 하둡 에코 시스템으로 구성되었지만 플랫폼 구성에 필요한 컴포넌트 중 하둡 에코 시스템에서 지원하지 않거나 기능이 미흡한 영역은 그루터 자체 기술을 이용하여 개발하여 플랫폼에 포함되어 있으며 수년간 빅데이터 시스템을 운영한 기술 및 지식이 집약된 시스템이다.

qoobah의 장점(1)

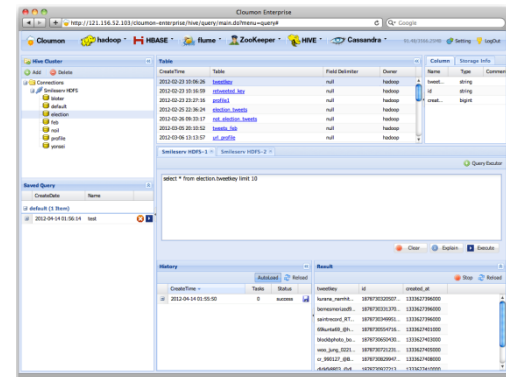
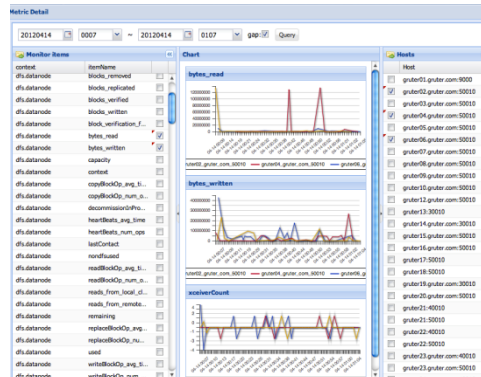
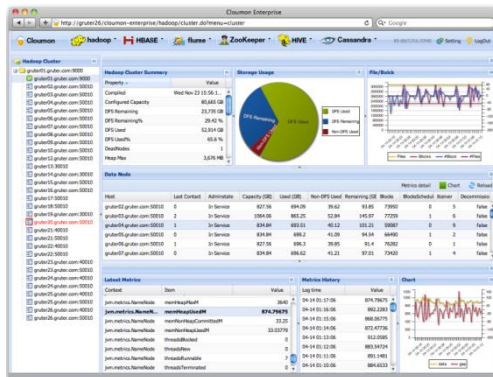
구분	오픈 소스 활용 시 문제점	qoobah를 이용한 해결 방안
아키텍팅	- 요구사항 대비 아키텍처 선정의 어려움 . 기반 기술의 복잡성 및 다양한 오픈 소스, 상용 솔루션 존재	-데이터 플랫폼 관련 높은 기술력 및 경험에 기반한 최적의 아키텍처 선정 . 수년간 자체 빅데이터 클러스터 운영 경험 . 오픈 소스 기반 빅데이터 프로젝트 다수 수행
	- 적절한 하둡 버전 선택 및 최적화의 어려움	- 사용자 요구 사항에 부합되는 하둡 버전 선택 . 1.0, CDH, Facebook release, 2.0 등
	- 요구사항에 적합한 오픈 소스 선정 문제 . FlumeOG vs. FlumeNG . 검색 플랫폼(katta, solr, elasticsearch) . NoSQL(HBase, Cassandra, MongoDB)	- 사용자 요구 사항에 적합한 오픈 소스 솔루션 선정 및 검증 . 솔루션 비교 검토 . 성능 및 안정성 테스트 등
기술 내재화	- 자체 개발 역량으로 오픈 소스에 대한 분석 및 노하우 축적 필요 . 장기간 숙련된 개발자 필요	- qoobah에 사용된 다양한 오픈 소스에 대한 교육 및 기술 전수 . 실제 운영 가능한 수준으로 기술 전수
데이터 수집	- 기본 제공 데이터 수집 모듈 기능 부족 . 필요 모듈을 직접 개발해서 사용해야 함	- 다양한 데이터 수집 모듈 제공 . Checkpoint tailer 등 다양한 데이터 수집 소스 모듈 제공 - 수집 대상 서버 Context 변경 상태 모니터링 및 알림 기능 . 주기적으로 수집 대상 서버에 웹 Context 추가 상태 모니터링 - 다양한 옵션을 제공하는 HDFS 저장 기능 . Rolling, flush 주기 설정 등 기능 제공 - Master 저장소와 Slave 저장소를 지정할 수 있는 Multi-Sink 기능 . 하나의 데이터 소스를 여러 용도로 활용
	- 데이터 수집 관리 시 직관적이지 않고 관리하기 어려운 UI	- 데이터 발생원으로부터 저장소까지 통합 관리 기능 제공 . 데이터 수집 대상 서버 관리 . 직관적인 Data flow 관리 . 웹 기반 수집 대상 추가 및 제거, 설정 변경 관리

qoobah의 장점(2)

구분	오픈 소스 활용 시 문제점	qoobah를 이용한 해결 방안
데이터 수집	- 데이터 수집 플랫폼 자체 취약점 존재 . Flume의 Multi-master 구성 시 configure 정보 미일치 문제	- Multi-master의 문제 해결 모듈 제공
	- 하나의 데이터 소스로 다양한 활용 시 유연한 구성 어려움	- 자체 개발한 데이터 버스 서비스인 Bamboo 제공 . 하나의 데이터 소스를 여러 용도로 사용할 수 있는 데이터 버스 기능 제공
	- 외부 데이터 수집 기능 미제공	- 자체 개발한 Grubot을 이용하여 블로그, 뉴스, 트위터 등을 수집 (qoobah 확장 모듈에서 제공)
데이터 분석 및 활용	- 실시간 분석 기능 취약 . 오픈 소스 라이브러리를 이용하여 직접 개발해야 함	- 실시간 분석 룰 관리 기능 기본 제공 . 분석 룰 관리 웹 UI 기능 . 분석 룰 데이터 수집기와 자동 연동 . 분석 결과 HBase 등에 자동 저장 및 검색 기능
	- 저장된 데이터 조회 기능 취약	- 비 개발자도 쉽게 접근할 수 있는 웹 기반 분석 기능 제공 . 웹 기반 Hadoop File Browser . Hive Workbench(Map/Reduce 관리 체계와 자동 연동)
	- 대용량 데이터 검색 기능 취약 및 비용 증가 . 로그 데이터는 볼륨의 문제로 과거 데이터 검색 시 비용 증가 . 검색 시스템 구성 어려움 발생	- 검색 응답 요구사항에 따라 다양한 검색 스토리지 구성 . 스토리지 티어링을 제공하는 자체 개발한 elasticsearch gateway 제공 . Local SAS/SSD, Local SATA, HDFS
	- 저장된 데이터를 이용한 실시간 분석 취약 . Storm, S4 등 추가 오픈 소스 모듈 채택 필요	- Hadoop 저장 파일을 큐로 이용하는 단순한 모델 기반의 자체 개발한 분석 플랫폼 제공
관리	- 관리 취약 . 각 오픈 소스별로 독립적인 관리 기능 제공 . 장애, 서버 상태 등에 따른 알람 체계 미지원	- 웹 UI 기반 통합 관리 기능 제공 . 하나의 관리 도구를 이용하여 다양한 오픈 소스 관리 . 단순 모니터링이 아닌 다양한 관리 기능 제공 . 서버 장애 등에 대해 알람 기능 제공 . 각 데몬의 다양한 Metrics 데이터를 과거 이력 정보까지 저장 . 장애 상황 발생 시 장애 원인 분석용으로 활용

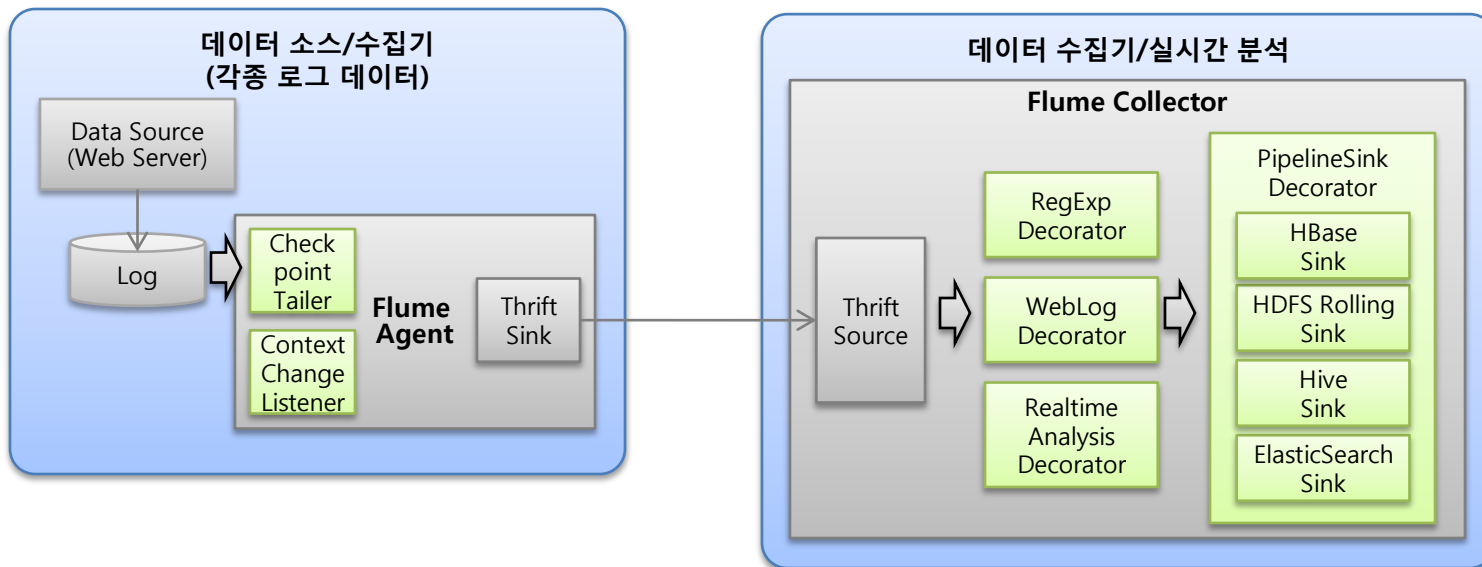
Cloumon

- Hadoop 기반 오픈 소스 통합 관리 및 모니터링
 - Hadoop File System: NameNode, DataNode 모니터링, 웹 기반 File browser
 - Hadoop MapReduce: JobTracker, TaskTracker 모니터링, MapReduce Job Management
 - Hive: Hive Query 실행, Hive Table Schema 관리, M/R 관리와 연동
 - ZooKeeper: ZooKeeper 데몬 모니터링, Znode browser 및 ACL, Watcher 관리
 - Flume: Flume Node 모니터링, Flume config 관리
 - HBase: Master, RegionServer 모니터링, Data/Schema 관리 work bench
 - Workflow: Oozie 기반의 workflow 생성 및 실행, 이력 관리
- 별도 솔루션으로 제공



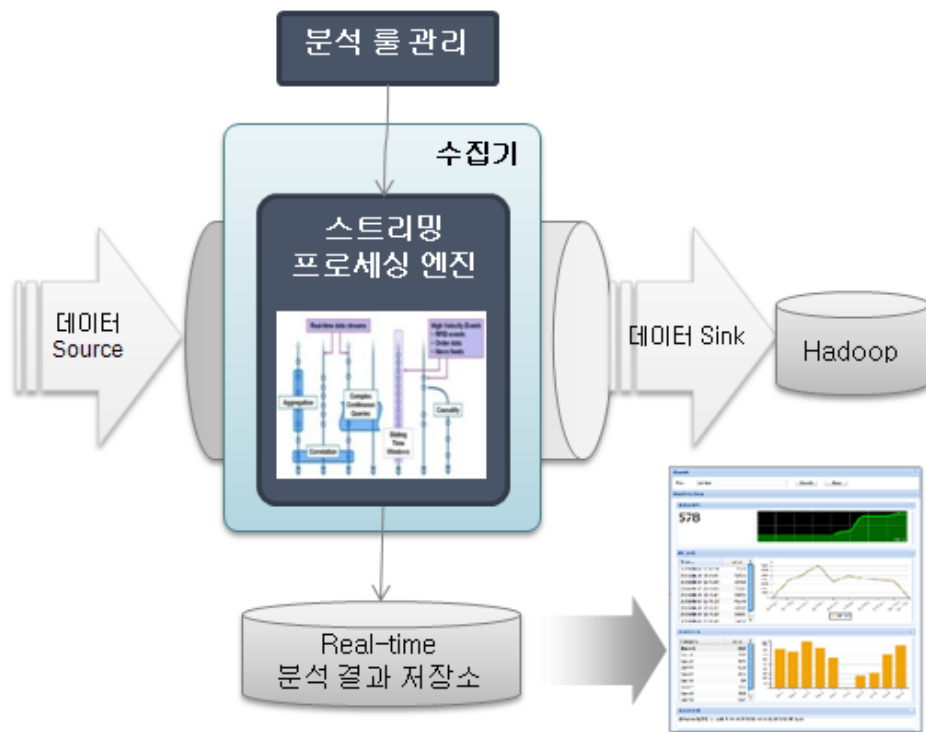
Flume plugin

- Flume 기본 제공 모듈에서의 취약점을 보완한 plug-in
 - 모니터링에 필요한 metrics 정보 제공(cloumon의 flume 관리와 연동)
 - Check point tailer: 전송 이력을 checkpoint 로 관리하여 agent 재시작 시 데이터 중복 전송 방지
 - Context change listener: Agent 설치된 서버에 새로운 웹 서버 설치나 Web Context의 변경 사항을 모니터링
 - RegExpDecorator: 로그를 Regexp 기반으로 파싱 처리
 - WebLog Decorator: 일반적인 표준 웹 로그에 대한 파싱 처리
 - Realtime Analysis Decorator: 동적을 변경되는 Esper 질의를 반영하여 처리할 수 있는 Sink
 - PipelineSinkDecorator: 하나의 로그 Event에 대해 여러 개의 Sink를 pipeline 으로 처리하게 하면서 master/slaves sink 지정
 - HDFS Rolling Sink: 시간 주기 또는 레코드 건수 단위로 Hadoop File을 rolling 하면서 저장하는 Sink
 - HBase/Hive/ElasticSearch Sink: HBase/Hive/ElasticSearch 등으로 로그 데이터를 저장하는 Sink



Realtime Analysis Query Manager

- Esper와 같은 실시간 분석 질의 관리
 - Flume의 Decorator와 Esper 연동
 - Time series, Incremental, Overwrite, Category 등의 분석 유형 제공
 - 분석 결과 저장은 HBase 이용(MySQL 등으로 교체 가능)
 - 기 정의된 분석 유형에 대한 분석 결과 조화 UI 제공

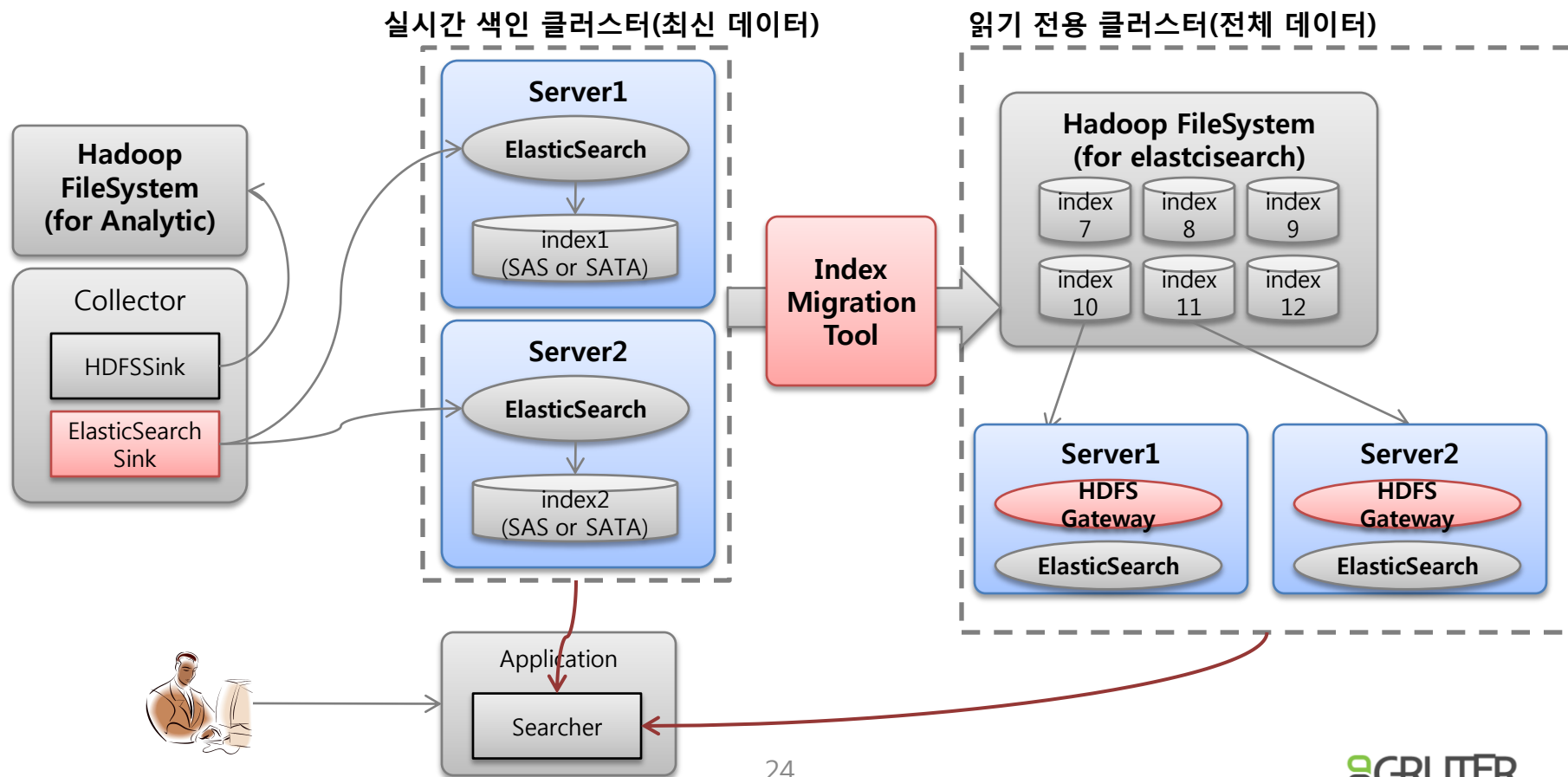


- 데이터 수집과 동시에 제한된 분석을 수행
- Collector에 내장되어 분석 수행(adaptor 역할)
- 분석 롤 관리시스템과 연동되어 분석 롤 관리
- 복잡한 분석보다 **count, sum** 등 단순한 **time window** 기반의 **aggregation** 분석 수행
- 분석 결과 데이터에 대한 정합성이 느슨한 업무에서 활용

- `select uid as user, 1 as cnt from TwitterEvent where text like '%olleh%'`
- `select uid, count(uid) from TwitterEvent.win:time_batch(60 sec) group by uid where text like '%olleh%'`

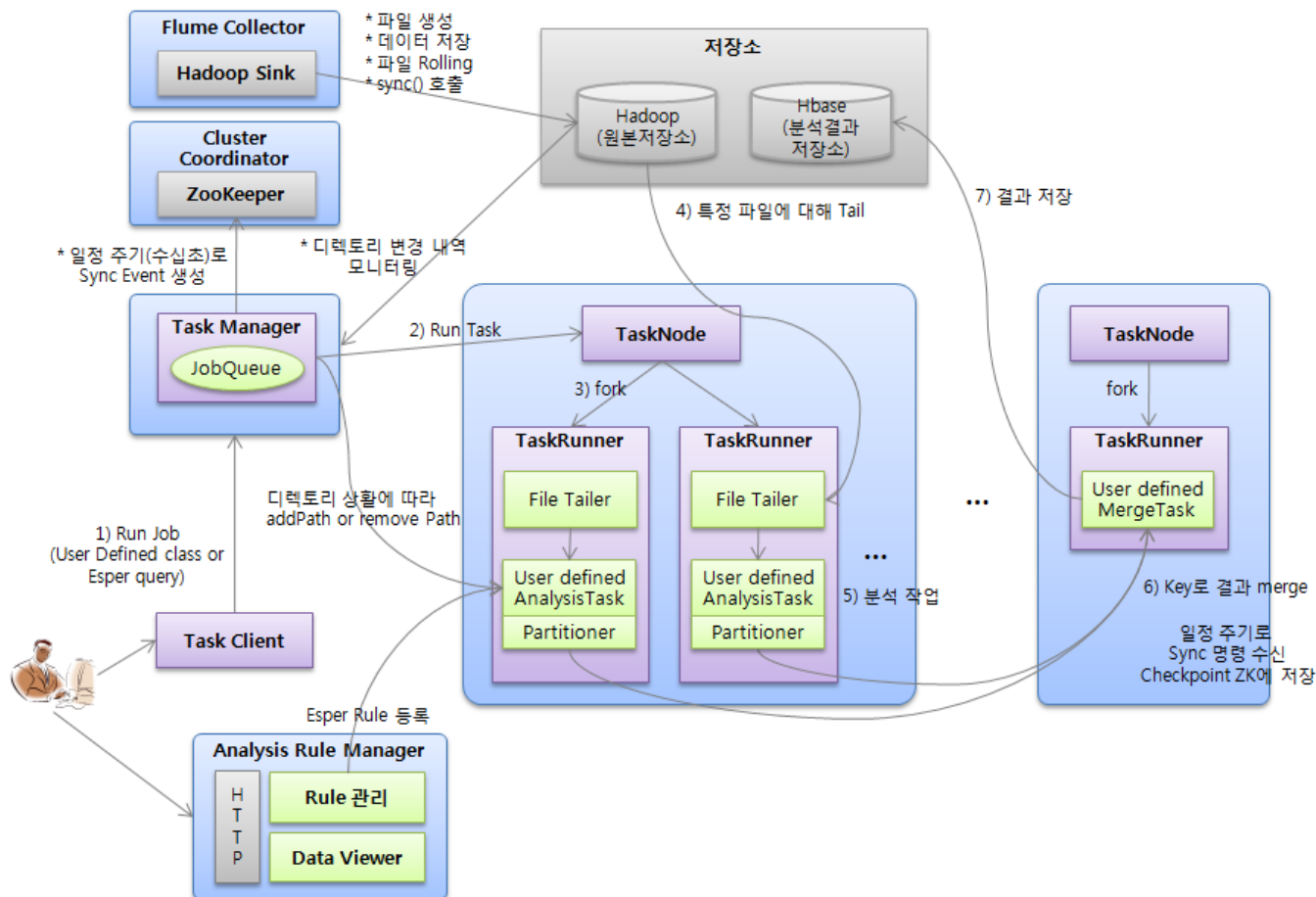
ElasticSearch Hadoop Plugin

- ElasticSearch 인덱스 저장소로 Hadoop File System을 이용하는 plug-in
 - 분산 검색 솔루션은 데이터 용량이 많아지면 실시간 인덱싱 및 검색에 부하 증가
 - 실시간 색인용 클러스터와 읽기 전용 클러스터 운영
 - 실시간 색인 클러스터는 기본 ElasticSearch 사용
 - 읽기 전용 클러스터는 ElasticSearch + Hadoop Plugin 사용



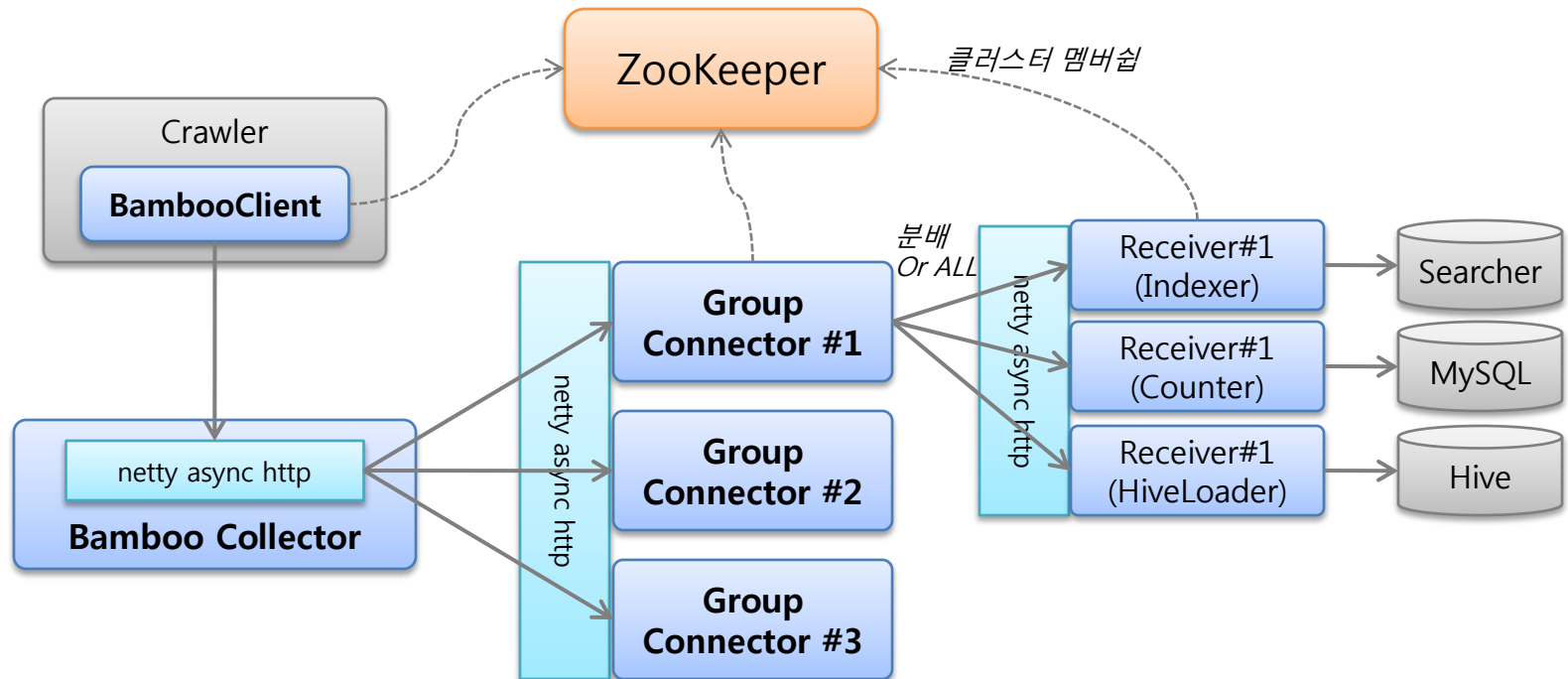
Cloustream

- HDFS에 저장된 데이터를 tail 하여 분석 모듈을 수행하는 준-실시간 분석 플랫폼
 - Hadoop 파일을 queue 데이터로 활용
 - 안정적인 checkpoint 관리로 컴퓨팅 노드 장애 시 재 시도 가능
 - 실시간 MapReduce의 개념으로 별도의 프로그램 모델 제공
 - Epser 등과 연동 가능

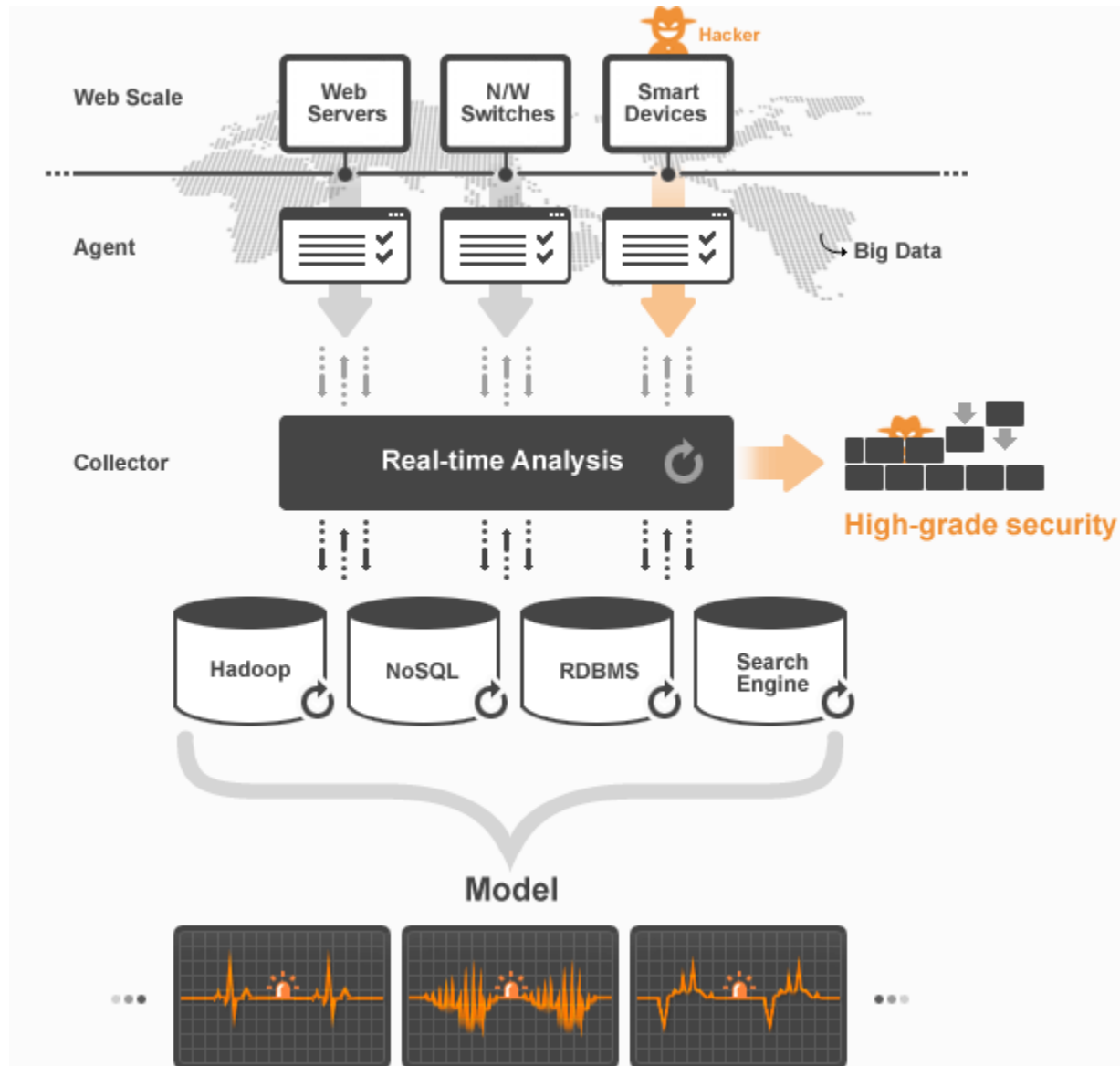


Bamboo

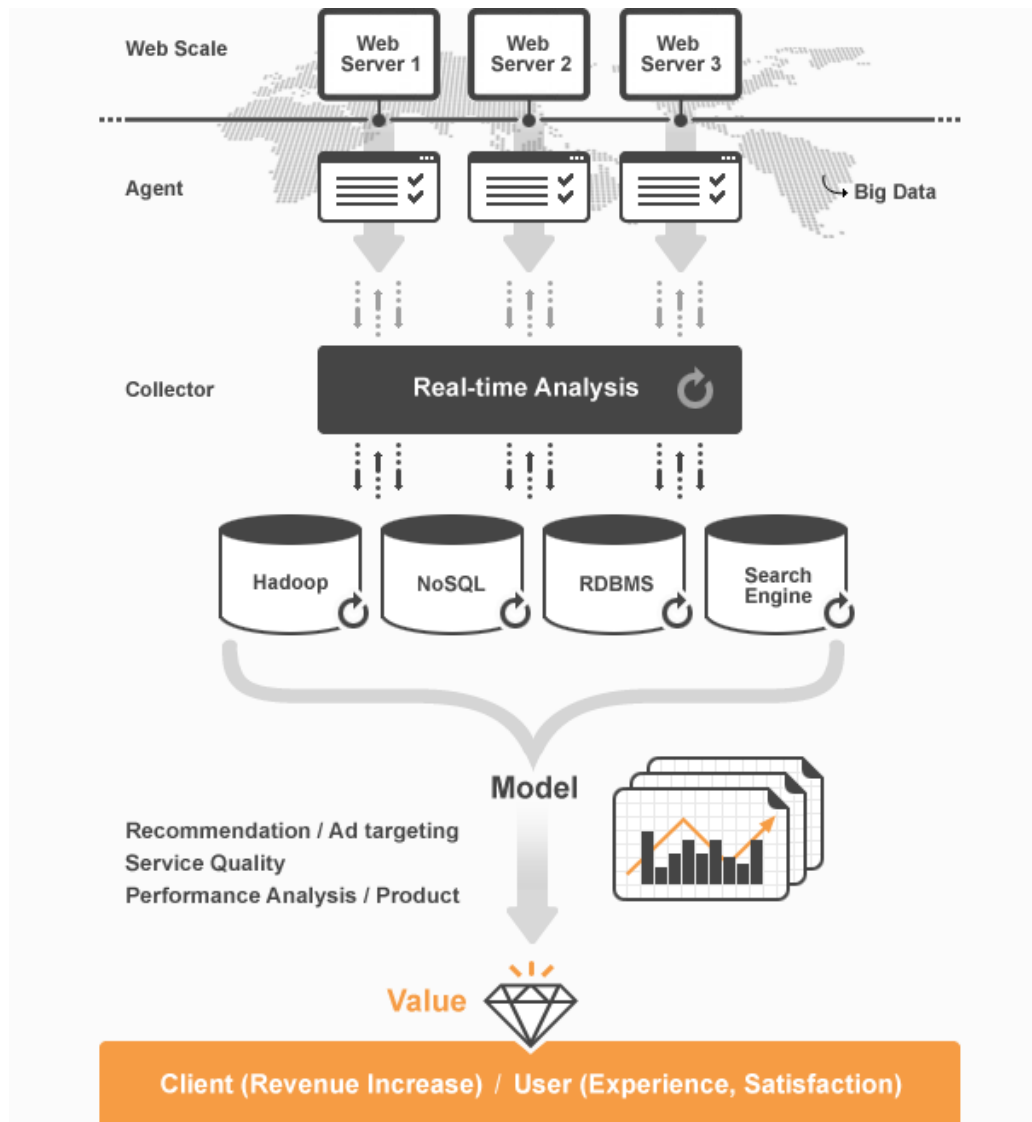
- 데이터의 스트림을 제어하면서 필요한 비즈니스 로직을 동적으로 추가할 수 있는 데이터 버스 플랫폼
 - 시스템 runtime 중에도 노드(Server/Client 조합) 연결을 통해 data flow 확장 가능 하고, 동적으로 프로세싱 모듈 조합/연결을 통해 data processing 확장 가능.
 - 각 노드는 upstream/downstream 구조(Flume의 source와 sink 개념과 유사)
 - 요구사항에 따라 flume으로만 구성할 것인지 bamboo와 결합해서 사용할 것인지 결정



Security log analysis system

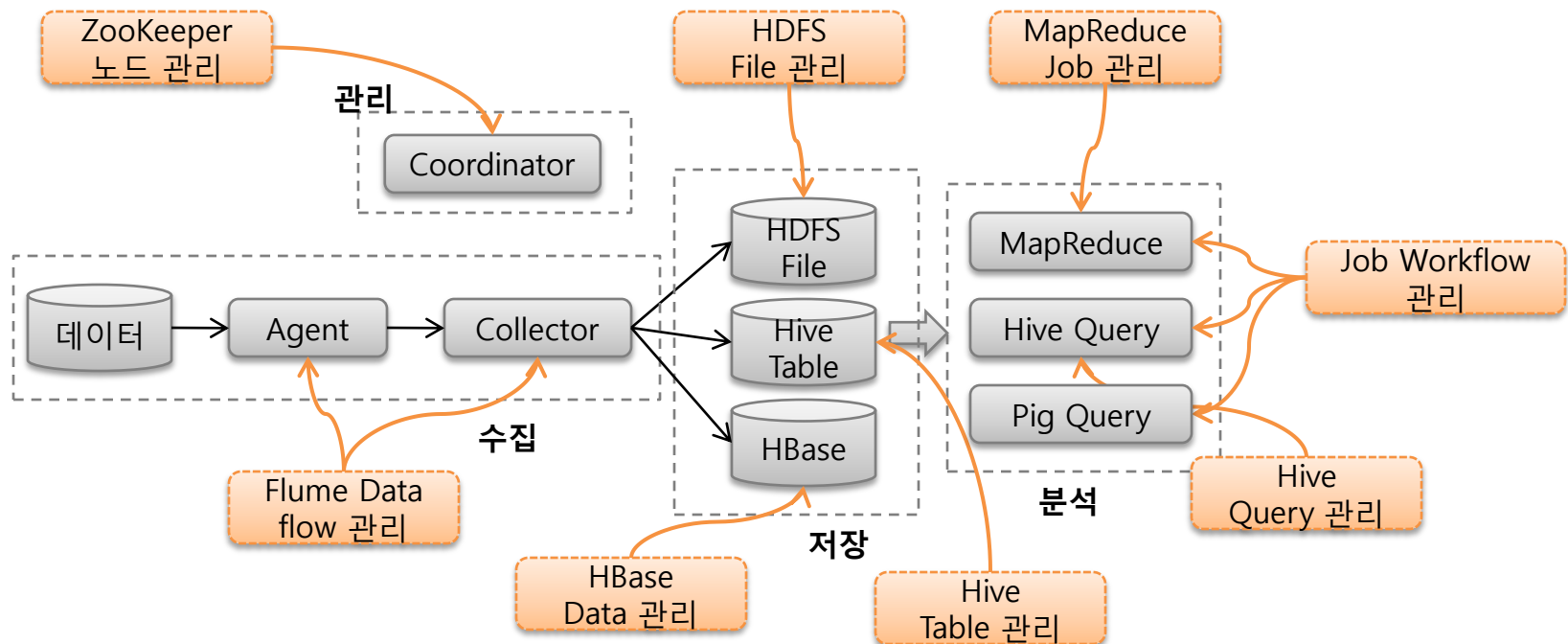


Real-time log analysis system

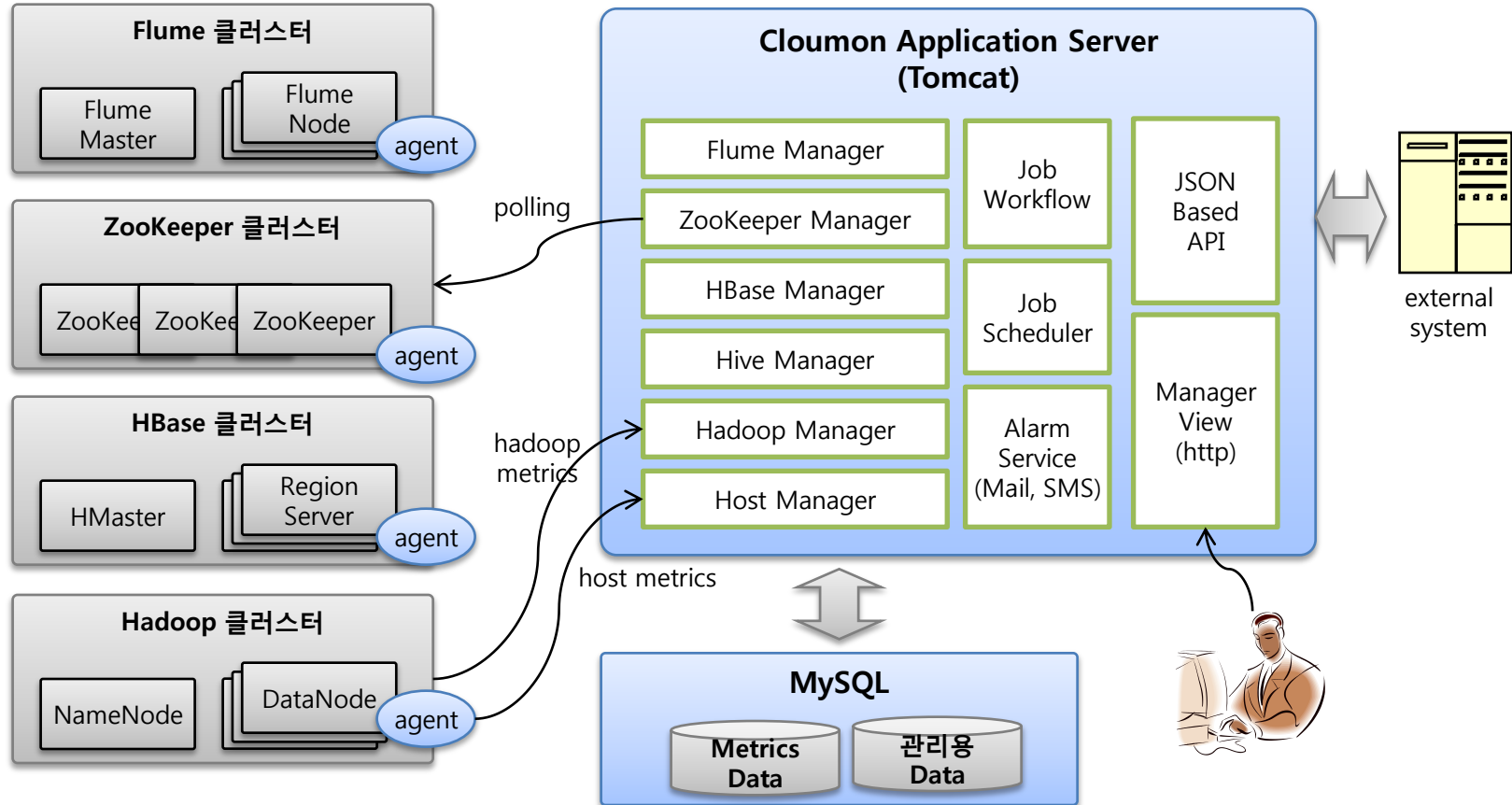


Cloumon

- Cloumon은 Hadoop과 Hadoop echo system을 구성하는 오픈 소스에 대한 관리, 모니터링을 기본적으로 제공
- 오픈 소스의 취약한 관리, 모니터링을 해결
- Hadoop 기반 빅데이터 플랫폼을 구성할 경우 플랫폼 전체를 유기적으로 관리하는 역할 수행

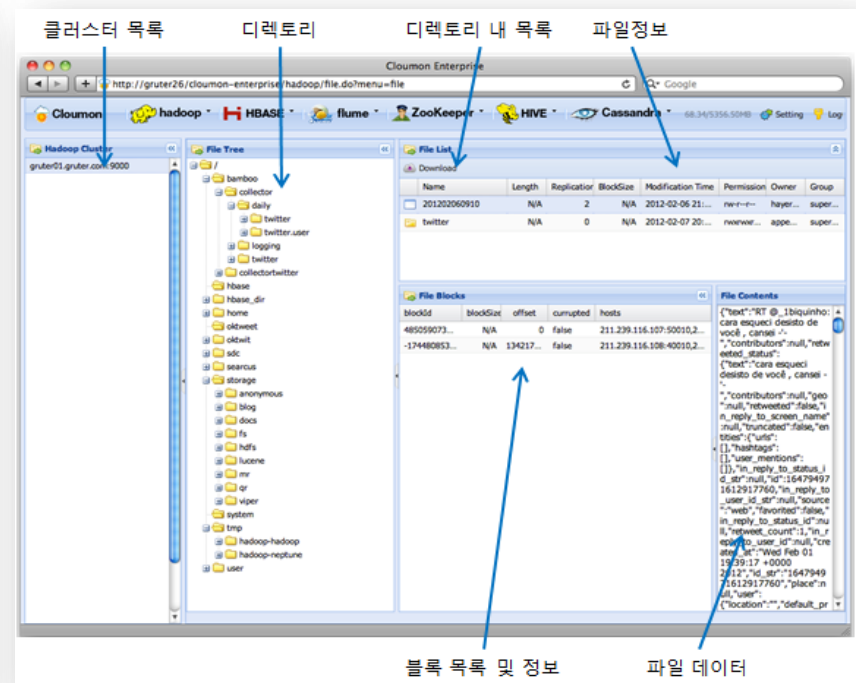
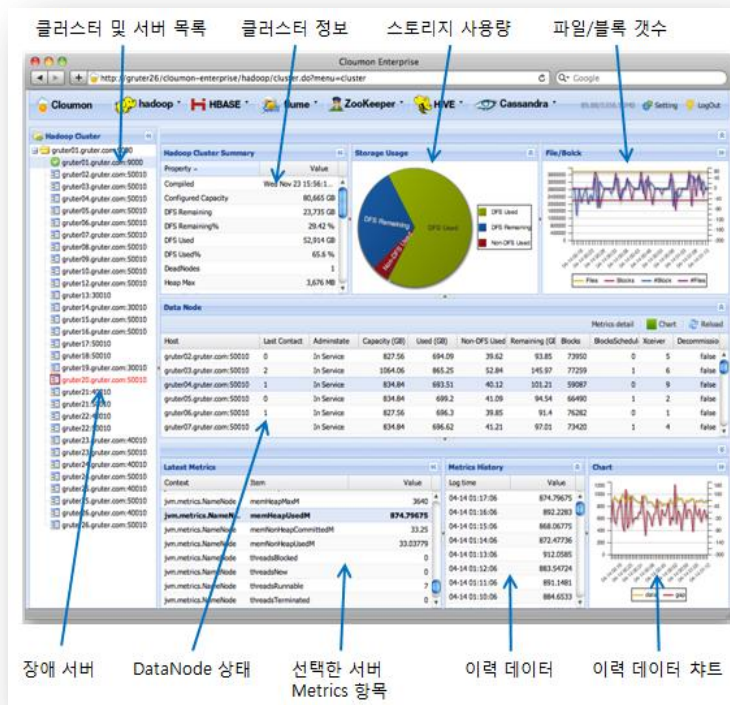


Cloumon: Architecture



Cloumon: HDFS

- 서버 모니터링 및 상태 정보 관리
 - NameNode, DataNode의 데몬 상태와 메모리 Heap 사용 상태, 스레드 개수 등을 관리
 - 비정상 상황 발생 시 알람
- 파일 관리
 - 웹 기반 파일 브라우저
- 멀티 클러스터 관리
 - 하나의 Cloumon으로 여러 HDFS 클러스터 관리 가능



Cloumon: MapReduce

- MapReduce 클러스터 관리
 - JobTracker, TaskTracker의 데몬 상태와 메모리 Heap 사용 상태, 쓰레드 개수 등 관리
- MapReduce Job 모니터링
 - MapReduce 작업 목록과 진행 상황을 관리
 - 작업의 Summary, Counter 정보, job.xml 정보 등 관리
 - Task Profiling

The screenshot shows the Cloumon web interface for monitoring MapReduce jobs. Annotations with arrows point to specific features:

- 클러스터 목록** (Cluster List): Points to the 'MapReduce Cluster' dropdown menu on the left.
- 작업 목록** (Job List): Points to the 'Jobs' table listing various jobs with their status and progress.
- 작업 상세 정보 선택** (Job Detail Selection): Points to the 'Job Detail' tabs (Detail, Job.xml, Map Task, Reduce Task).
- Map/Reduce 진행상태** (Map/Reduce Progress): Points to the 'Task summary' table showing the progress of Map and Reduce tasks.
- Job Counter**: Points to the 'Counter' table showing various counters like 'Launched reduce tasks' and 'Data-local map tasks'.

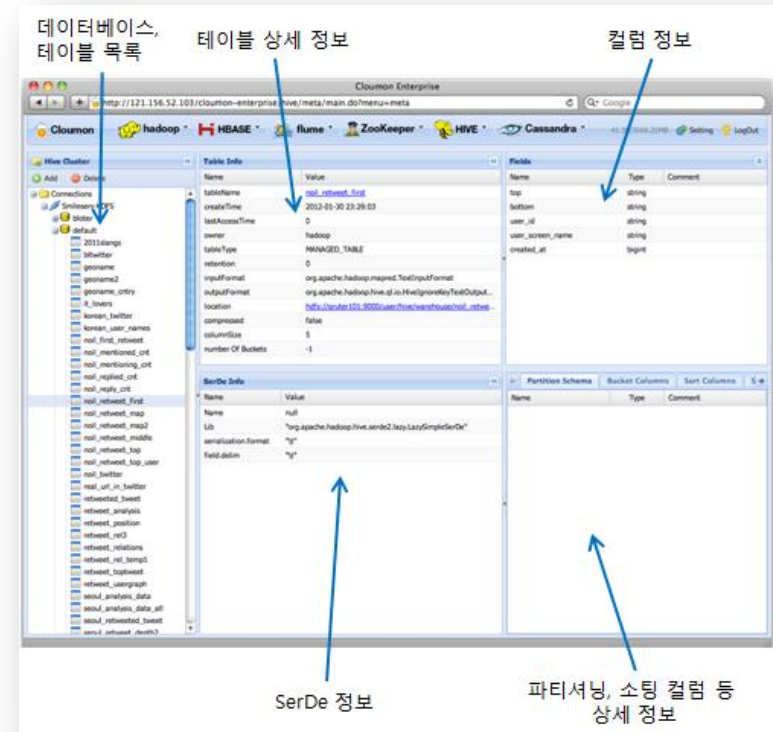
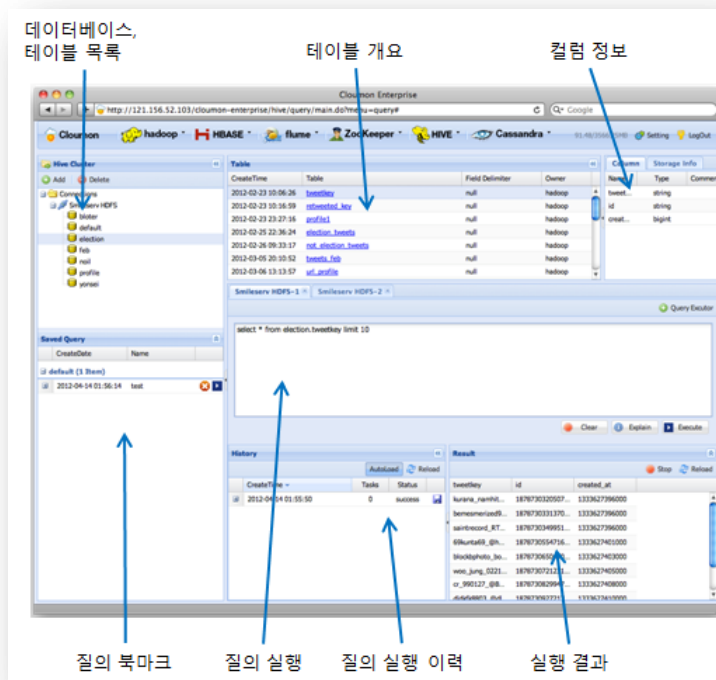
Jobid	status	User	Name	Map(%)	Maps	Map Completed	Reduce(%)	Reducers	Reduce Complete	Schedule
job_201204252030_0005	success	sel...	sel...	100.00%	60	60	0.00%	0	0	
job_201204252030_0006	success	sel...	sel...	100.00%	4	4	0.00%	0	0	
job_201204252030_0007	running	im...	im...	91.63%	6480	5986	0.91%	999	0	
job_201204252030_0008	running	sel...	sel...	0.00%	4	0	0.00%	0	0	

Kind	% Complete	Num Tasks	Pending	Running	Complete	Killed	Failed
Map	91.10%	6480	539	30	5991	0	0
Reduce	0.91%	999	969	30	0	0	0

Group Name	Item Name	Map	Reduce	Total
Job Counters	Launched reduce tasks	0	0	0
	Local map tasks	0	0	4,394
	Launched map tasks	0	0	5,921
	Data-local map tasks	0	0	1,527
org.apache.hadoop.hive...	PASSED	2,541,905...	0	2,541,905,939
	FILTERED	3,566,124...	0	3,566,124,519
FileSystemCounters	FILE_BYTES_READ	399,805,61...	0	399,805,618,438

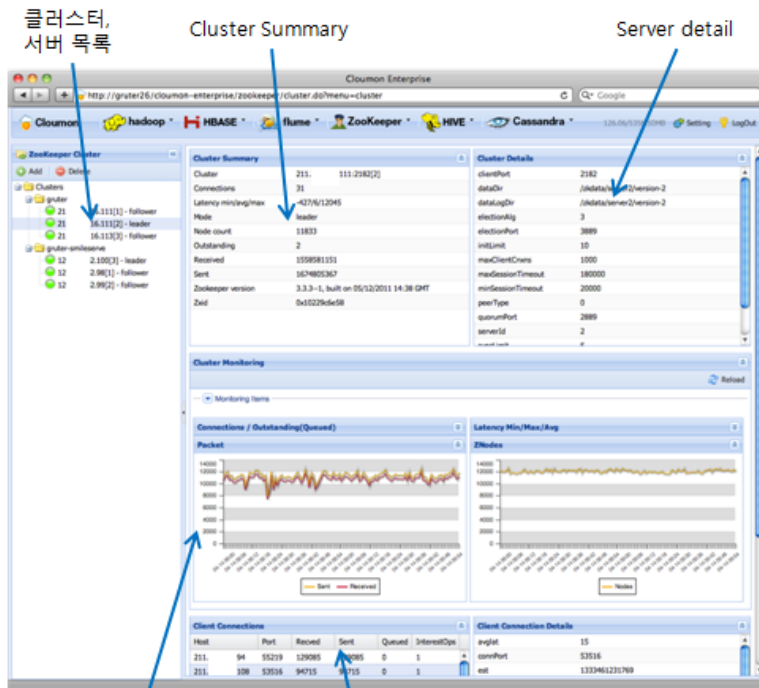
Cloumon: Hive

- 스키마 관리
 - Hive의 데이터베이스, 테이블 정보를 관리하는 기능
 - 테이블 목록, 생성, 삭제 관리
 - 테이블의 컬럼 정보, SerDe 정보, 파티션 정보 등 상세 정보 제공
- 질의 Workbench
 - HiveQL을 실행하고 실행 결과를 브라우저에서 관리할 수 있는 기능
- MapReduce 작업 모니터링 연동



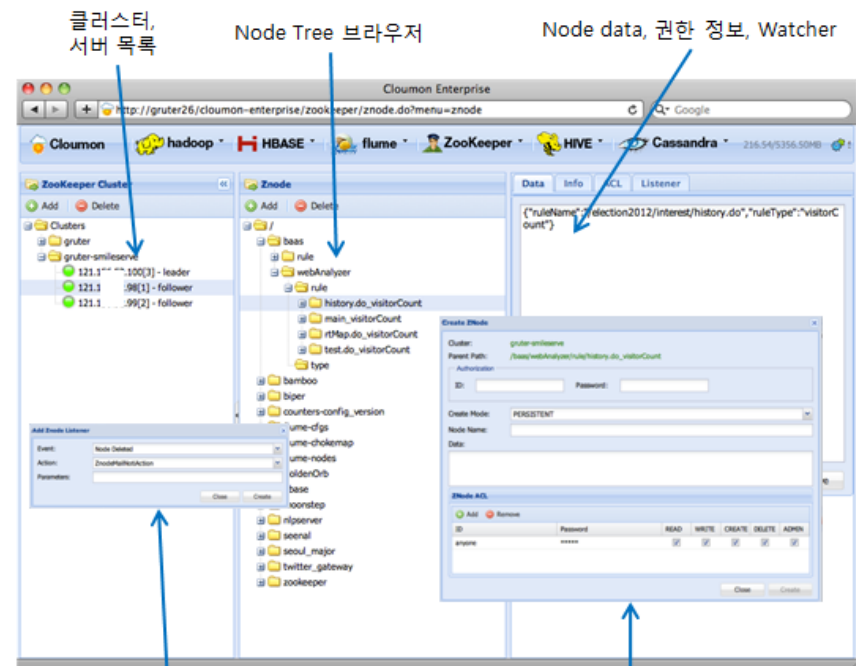
Cloumon: ZooKeeper

- ZooKeeper 클러스터 관리
 - ZooKeeper 클러스터를 구성하는 서버의 상태 정보 및 각 서버에 접속되어 있는 클라이언트의 접속 상태와 네트워크 트래픽, 큐 대기 시간 등을 모니터링 할 수 있는 기능
- 노드 관리
 - 탐색기와 같은 구조로 노드의 구조를 보여주고 노드를 쉽게 추가, 삭제, 수정할 수 있는 기능
 - 노드 권한 관리
 - 노드 Watcher 등록



Queue, Latency,
Packet, node 수

접속된 Client



Watcher 등록

노드 생성

Cloumon: Flume

- Flume 클러스터 관리/모니터링
- Data Flow 관리
- 스마트한 노드 관리
- 편리한 config 생성 툴 제공
- 다양한 Source, Sink, Decorator 제공

The screenshot displays the Cloumon Enterprise web interface. At the top, there's a navigation bar with icons for Cloumon, Hadoop, HBase, Flume, ZooKeeper, Hive, and Cassandra. The main area is divided into several panels. On the left, a 'Clusters, Server List' panel shows a list of nodes. The central part features a 'Data Flow (Agent)' and 'Data Flow (Collector)' section with a visual flow diagram. Below this is a 'Flume Config Manager' section with tabs for 'Physical Nodes' and 'Logical Nodes'. The 'Logical Nodes' tab is active, showing a table of nodes with columns for name, physical id, status, flow id, source, and sink. At the bottom, there's a 'Flow 맵핑' (Flow Mapping) section with a visual representation of the data flow.

클러스터, 서버 목록

Data Flow(Agent)

Data Flow(Collector)

Flow 맵핑

Config 디자이너

The screenshot displays the Cloumon Enterprise web interface, focusing on the 'Node Status' and 'Node Configuration History' sections. The 'Node Status' section shows a table of nodes with columns for Host, LogicalNode, PhysicalNode, Status, Version, and Lastseen. The 'Node Configuration History' section shows a table of configurations with columns for Node, Version, Flow ID, Source, and Sink.

클러스터, 서버 목록

서버 상태 정보

Config History

Host	LogicalNode	PhysicalNode	Status	Version	Lastseen	De	Lastseen
gruter103.gruter.com	gruter103-agent	gruter103-agent	ACTL	2012-04-09 12...	2	2012-04-14 2...	
gruter105.gruter.com	gruter105-coll...	gruter105-coll...	ACTL	2012-04-09 12...	1	2012-04-14 2...	
gruter106.gruter.com	gruter106-web...	gruter106-web...	IDLE	None	2	2012-04-14 2...	
seemal11.gruter.com	seemal11-web...	seemal11-web...	DST	2012-04-09 21...	422301	2012-04-10 0...	

Node	Version	Flow ID	Source	Sink
agent	2012-04-09 12:34...	default-flow	bambooSource()	agentSink("gruter106",8088)
collector	2012-04-09 12:34...	default-flow	thriftSource(8088)	multi("bambooHdfs","hdfs://grul...
gruter103-agent	2012-04-09 12:34...	default-flow	bambooSource	agentSink("gruter105",8088)
gruter105-collector	2012-04-09 12:34...	default-flow	thriftSource(8088)	multi("bambooHdfs","hdfs://grul...
seemal11-webAnalyzer	2012-04-09 21:28...	default-flow	tail("user/local/httplog/access...	(webAnalyzer("gruter101:2181.g...

Cloumon: Workflow

- Workflow application 디자인
- Workflow 작업 스케줄링 및 이력 관리

App 목록 Job 상세 설정 Workflow 디자인 Workflow App 상세 정보

The screenshot displays the Cloumon Enterprise web interface. On the left, the 'App List' table shows applications like 'hive_test', 'java_app', 'ssh_test', and 'test01'. The 'App Design' section shows a workflow diagram with nodes 'start', 'hive_test', 'pig-test', and 'end'. The 'Application property' panel on the right shows details for 'hive_test', including its creator, description, and XML definition. The 'Job Conf' section at the bottom allows for job configuration, including job type, name, and streaming options.

App 목록 Scheduled Job 목록 작업 상세 정보 작업 실행 이력

This screenshot shows the 'Scheduled Jobs' and 'Job Execution History' sections of the Cloumon Enterprise interface. The 'Scheduled Jobs' table lists jobs like 'test_job' with their schedules and last execution times. The 'Job Execution History' table provides a detailed log of job executions, including job names, app names, statuses, and timestamps. The 'Job Details' panel on the right shows the configuration for a specific job, including its name, app path, and execution status.

Cloumon: HBase

- HBase 클러스터 모니터링
- HBase Data 관리

RegionServer 목록 RegionServer 상태 RegionServer Metrics Detail

The screenshot displays the Cloumon HBase management interface. On the left, a tree view shows the 'Hbase Cluster' with a list of RegionServers. The main panel is divided into three sections: 'Region Servers' (listing online servers with their hostnames, start codes, and loads), 'Master Attribute' (showing cluster configuration like HBase Version and Zookeeper Quorum), and 'User Tables' (listing tables like 'RT_ANALYSIS'). Blue arrows point from the labels above to the corresponding sections in the interface.

Online	Host	Start Code	Load	Info Link
1	gruter101.gruter.com:60020	1314333035879	requests=0, regions=5, usedHeap=37, maxHeap=987	121.156.52.101.60030
2	gruter103.gruter.com:60020	1314333035538	requests=0, regions=5, usedHeap=35, maxHeap=987	121.156.52.101.60030
3	gruter104.gruter.com:60020	1314333026148	requests=0, regions=5, usedHeap=33, maxHeap=987	121.156.52.101.60030
4	gruter105.gruter.com:60020	1314333037282	requests=0, regions=5, usedHeap=25, maxHeap=987	121.156.52.102.60020
5	gruter102.gruter.com:60020	1314333755522		gruter102.gruter.com:60020,13143...

Attribute	Value
Hbase Version	0.90.3.1100300
Hbase Compiled	Sat May 7 13:31:12 PDT 2011,at...
Hbase Root Dir...	hdfs://gruter101:9000/hbase
Load average	5
Zookeeper Quo...	gruter103-2181.gruter102-2181...
Hadoop Version	0.20.3-GNAPSHOT
Hadoop Compl...	Fri Aug 5 21:04:50 KST 2011,had...

Table	Description
ORIGIN_DATA	(NAME => 'ORIGIN_DATA', FAMILIES => [(NAME => 'DATA1', BLOOMFILTER => ...
RT_ANALYSIS	(NAME => 'RT_ANALYSIS', FAMILIES => [(NAME => 'DATA1', BLOOMFILTER => ...

The screenshot displays the Cloumon HBase management interface. The top navigation bar includes links for Cloumon, Hadoop, HIVE, OOZIE, Flume, ZooKeeper, and HBASE. The main panel is divided into two sections: 'HBase Data View' (showing a list of tables with their names, authors, and links) and 'HBase Table View' (showing the schema of a selected table, including column names and data types). The 'HBase Data View' section is currently active, showing a list of tables with columns for 'rowKey', 'author', and 'link'.

rowKey	author	link
ac.kioot.blog/90149344103	author	http://blog.kioot.ac
ac.kioot.blog/90149344104	author	http://blog.kioot.ac
ac.kioot.blog/90149344105	author	http://blog.kioot.ac
ac.kioot.blog/90149344106	author	http://blog.kioot.ac
ac.kioot.blog/90149344107	author	http://blog.kioot.ac
ac.kioot.blog/90149344108	author	http://blog.kioot.ac
ac.kioot.blog/90149344109	author	http://blog.kioot.ac
ac.kioot.blog/90149344110	author	http://blog.kioot.ac

Cloumon: 마일스톤

- 2012.10.30: ver 2.0
 - Oozie 기반 workflow
 - HBase Data workbench
 - Esper QL Management
- 2012.02.28: ver3.0
 - Hadoop 2.0 지원
 - Mahout 등 분석 플랫폼 탑재
 - 그래프 생성 및 관리 기능 제공

DEMO

데이터 분석 서비스

Gruter's Hadoop History

- Hadoop 이전
 - 데이터 저장 서버 5대 수집 서버
 - Windows Server, Oracle, Local Disk
 - 수집된 블로그 데이터는 Oracle에 저장
 - 서버 장애 시 수천만 데이터 유실 및 재 수집
 - 수집 및 검색 서버 5대
 - Windows Server
- 2007.04 Hadoop 클러스터 최초 구성
 - x86 서버 10대, CentOS
 - Hadoop 0.9.0 설치
- 2008.05.05 NameNode 장애
 - Rack 전원 장애
- 2009 상반기
 - 10대 추가
 - Cloudata 10대 설치 및 운영
- 2009 NameNode 장애
 - NameNode Image Disk Full
- 2009 하반기
 - 6대 추가
- 2011 상반기
 - 7대 추가, 별도 IDC 구성
 - Hive 기반 분석
- 2012
 - NameNode 교체(32bits -> 64bits)
 - Hadoop 1.0.3 사용

Seenal
http://www.seenal.com/
Google

seenal
아이디
비밀번호
로그인

seenal^{beta}에서 소셜 미디어에 참여하고 성과를 측정하세요.

쉽고 상세한 수치로 성과 평가

부정, 분류별 메시지

대화 타겟 지정

관련어를 통해 메시지 한 눈에 보기

피트린 메시지로 중요 메시지 파악

지금, 신청하세요

현재 seenal은 클로즈 베타 테스트를 진행하고 있습니다. 베타 테스트가 되어보세요!

seenal 클로즈 베타 신청 >

좋아요

친구 중 제일 먼저 "좋아요"를 클릭하세요.

트윗 0

팔로우 @seejunajul

1 편리한 기능

아무리 수준 높은 분석이 돼도, 복잡하고 어려우면 바로 사용할 수 없습니다. 기업 담당자들의 이야기를 듣고, 실무자의 눈높이에서 개발했습니다.

2 확장성 높은 분류

기업에게 쏟아지는 메시지를 한눈에 알아 보기는 어렵습니다. 다양하게 분류해서 분류별로 메시지를 나누어 볼 수 있도록 했습니다.

3 수치화된 ROI 측정

기업의 대부분의 활동은 수량화 되어 측정하고 평가 됩니다. 업무 프로세스에 반영할 수 있도록 다양하게 수치화된 지표를 개발해 놓았습니다.

4 실시간 데이터 분석

소셜 미디어에서 고객과 소통하는 일은, 생각을 다루는 일입니다. 실시간으로 고객의 목소리를 듣고, 즉각 성과를 측정하고 대응할 수 있도록 개발했습니다.

5 보고서 커스터마이징

복잡한 엑셀작업으로 보고서를 작성할 필요 없습니다. 원하는 지표들을 정리해 보고서 템플릿을 만들어 놓으면 메일로 보내기, 파일로 보내기 등을 쉽게 할 수 있습니다.

6 메시지 확산 추적

소셜미디어에서 메시지가 어떻게 확산되는지 추적하기는 여간 어려운 일이 아닙니다. 모든 메시지에 대해 누가 누구에게 전파하는지 한눈에 확인할 수 있습니다.

seenal
© Copyright 2012 Gruter, Inc. All rights reserved.

저장/배치 분석 플랫폼 만들기

데이터 처리에 있어서 가장 큰 고민 중에 하나는, 미래의 구체적인 요구 사항을 아직 모른다, 다만, 확실한 것은:

- 데이터는 늘어날 것이다
- 데이터의 소비 용도도 다양해 질 것이다
- 데이터 프로세싱에 대한 다양한 요구도 늘어날 것이다

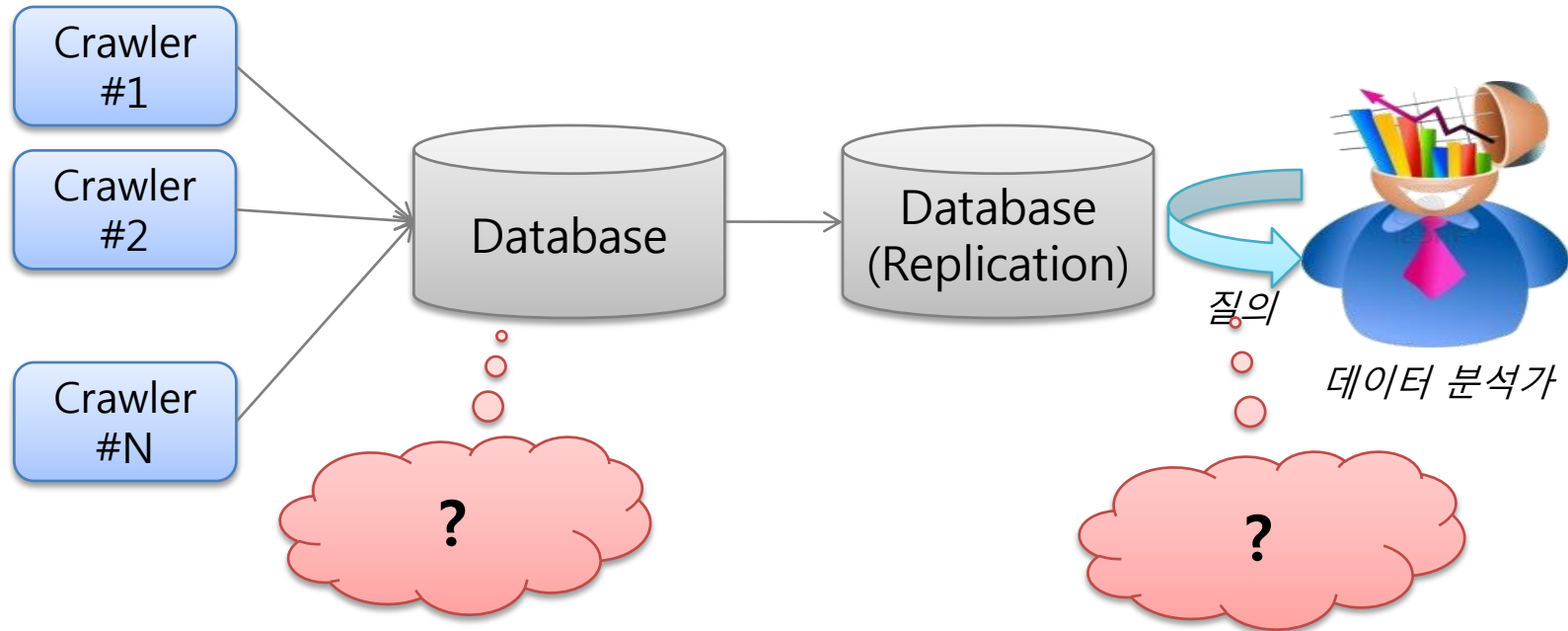
데이터의 흐름은 한번 시작되면 멈추지 않는다; 즉 달리는 차를 멈추고, 바퀴를 바꿔야 하는 식의 아키텍처는 맞지 않다.

기존에 구축된 시스템의 데이터 흐름에 영향을 없게 하거나 최소화 하면서 확장 요건을 만족 시키는 솔루션이 필요하다.

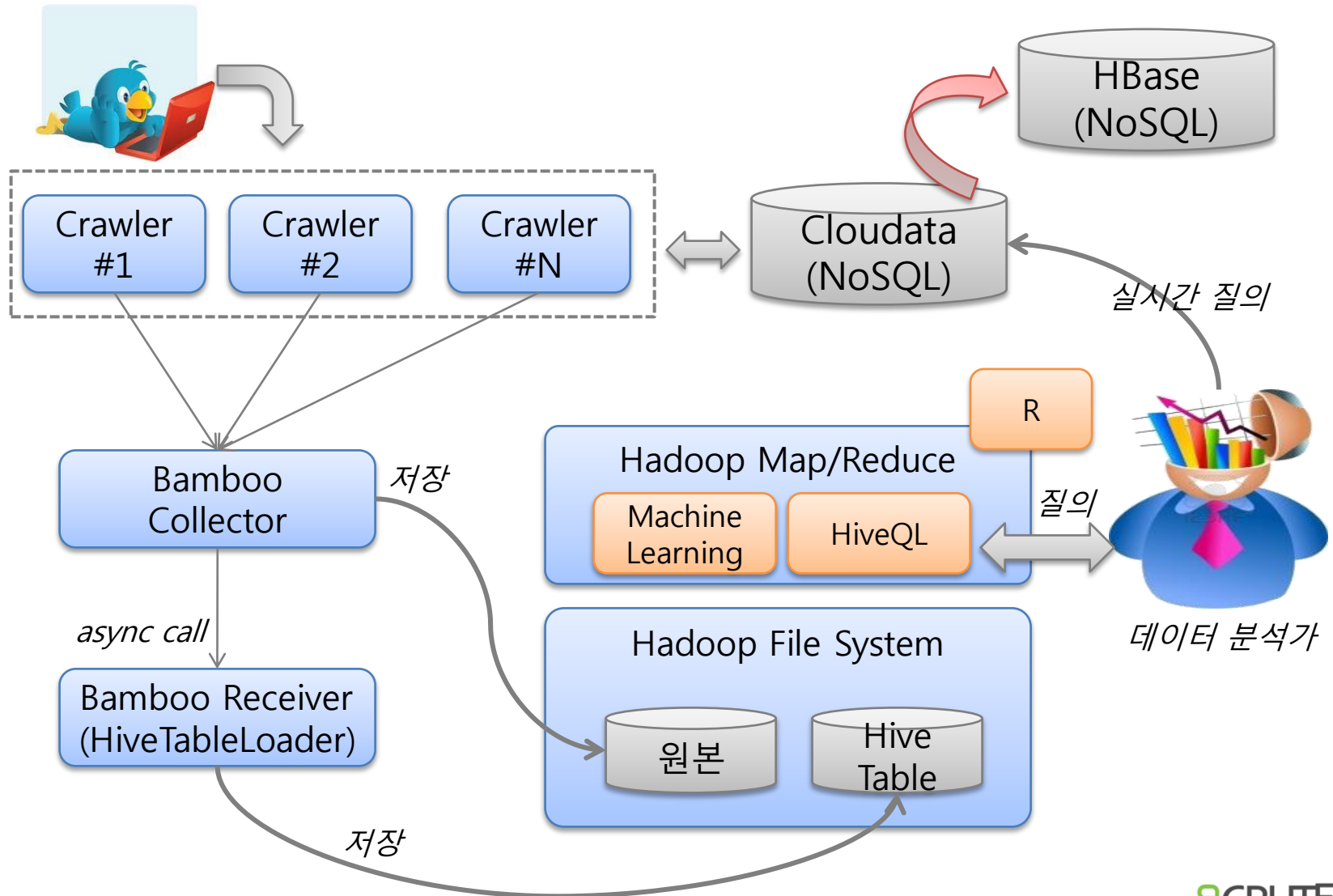


장정식
그루터
데이터 아키텍트

이런 구성은?

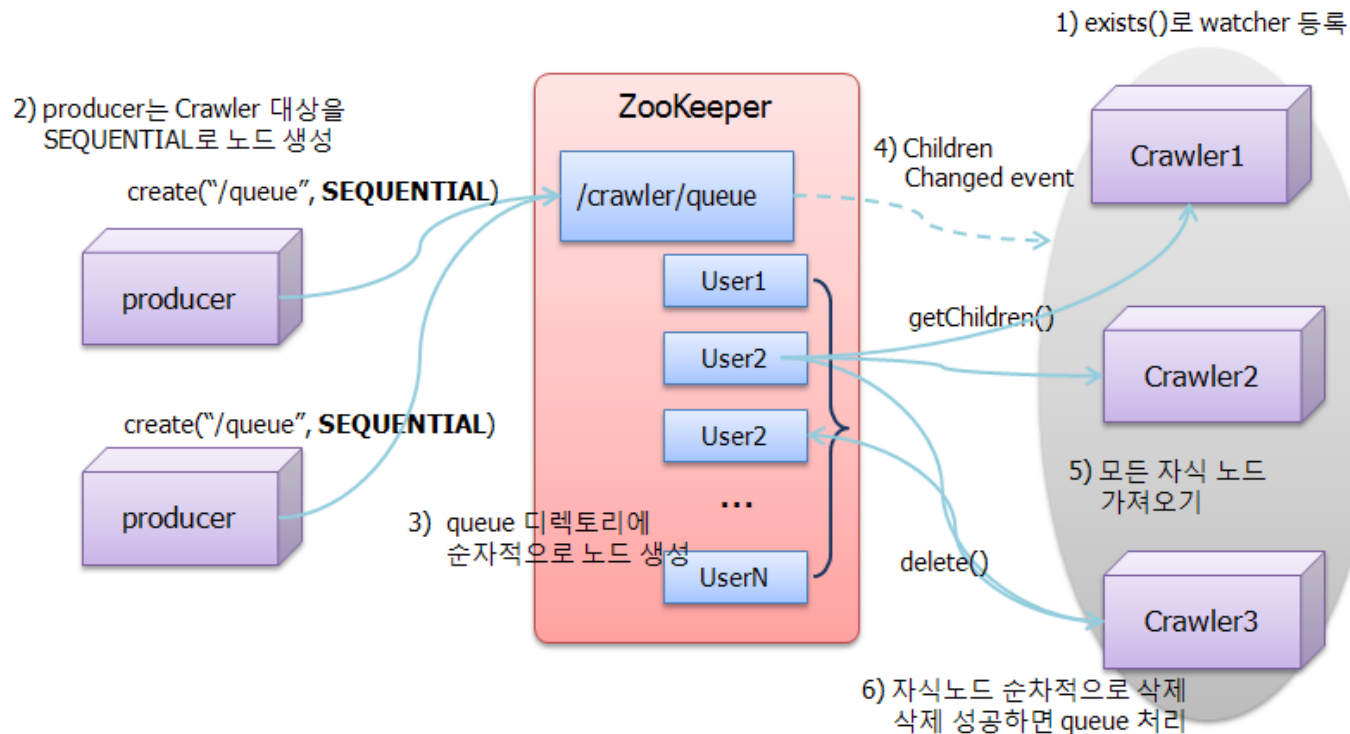


저장/배치 분석 플랫폼 구성



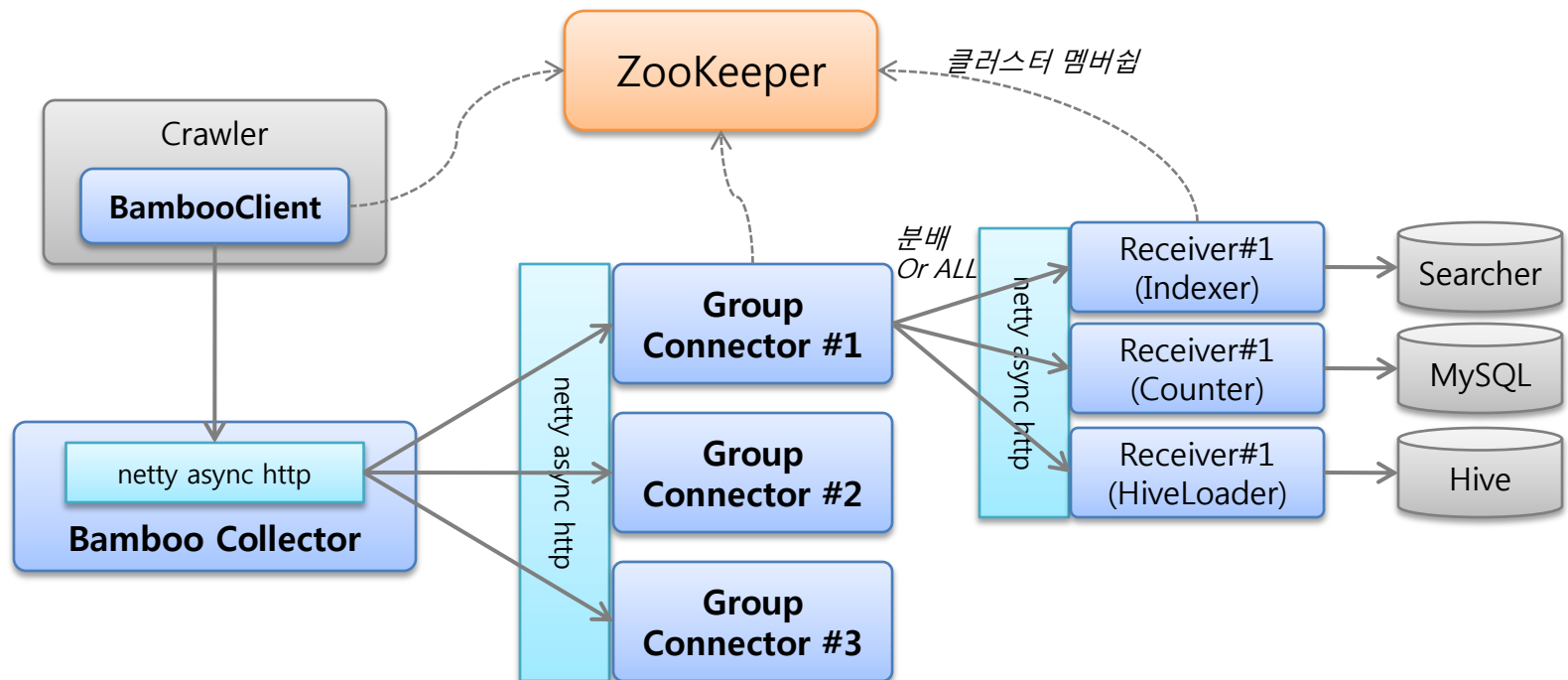
데이터 수집

수집 대상, 수집 데이터 증가에도 프로세스 증설만으로 수집 능력 향상
특정 Crawler 장애 시 자동으로 다른 Crawler가 역할 대신 수행

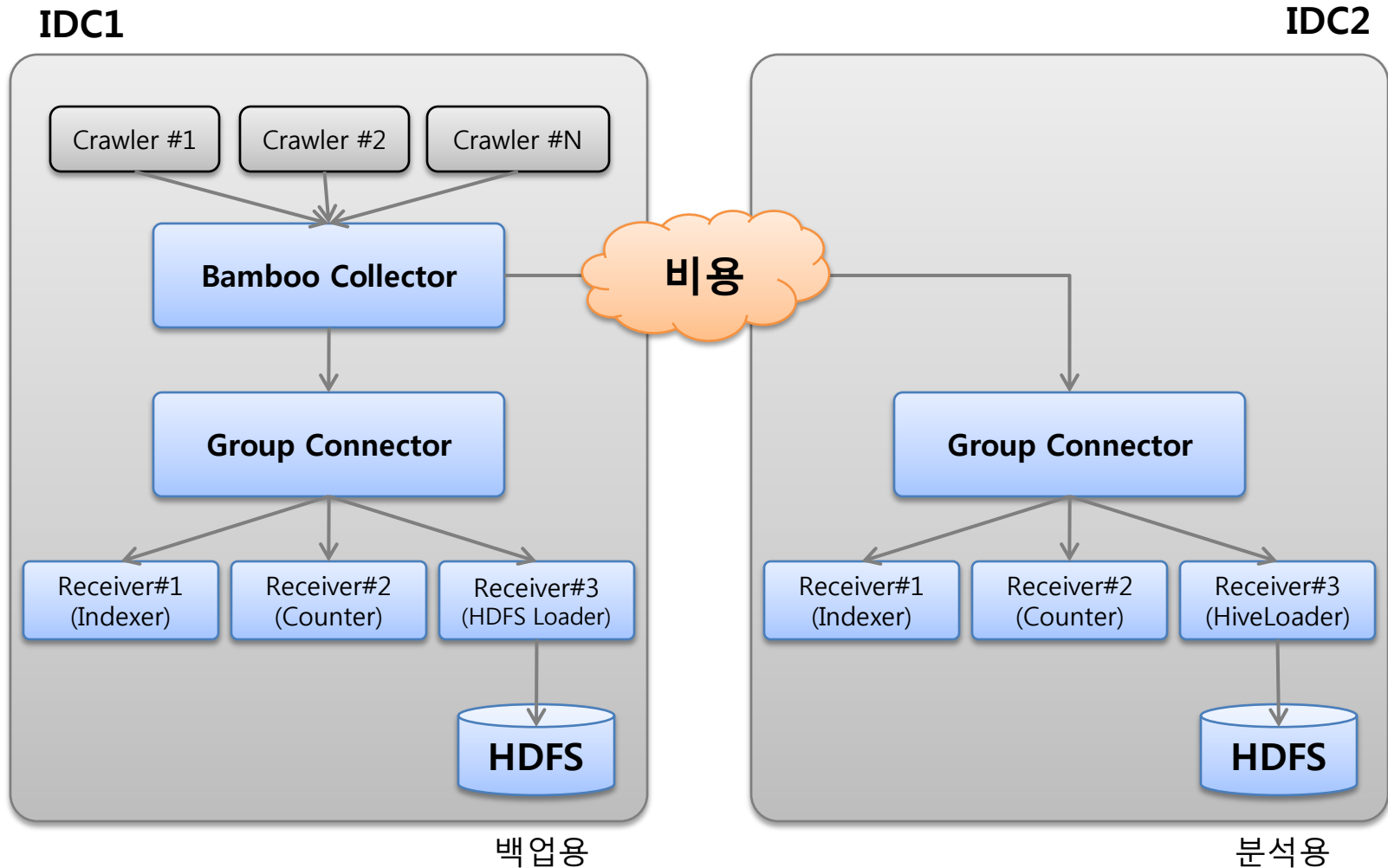


Bamboo: Gruter's data stream platform

- 데이터 발생원으로부터 데이터 처리와 흐름제어를 통해 목적지까지 수집된 데이터를 효과적으로 전달
- 각 노드는 Netty 기반의 upstream/downstream 구조(Flume의 source와 sink 개념과 유사)
- 시스템 runtime 중에도 노드(Server/Client 조합) 연결을 통해 data flow 확장 가능하고, 동적으로 프로세싱 모듈 조합/연결을 통해 data processing 확장 가능.



Bamboo를 이용한 IDC간 미러링



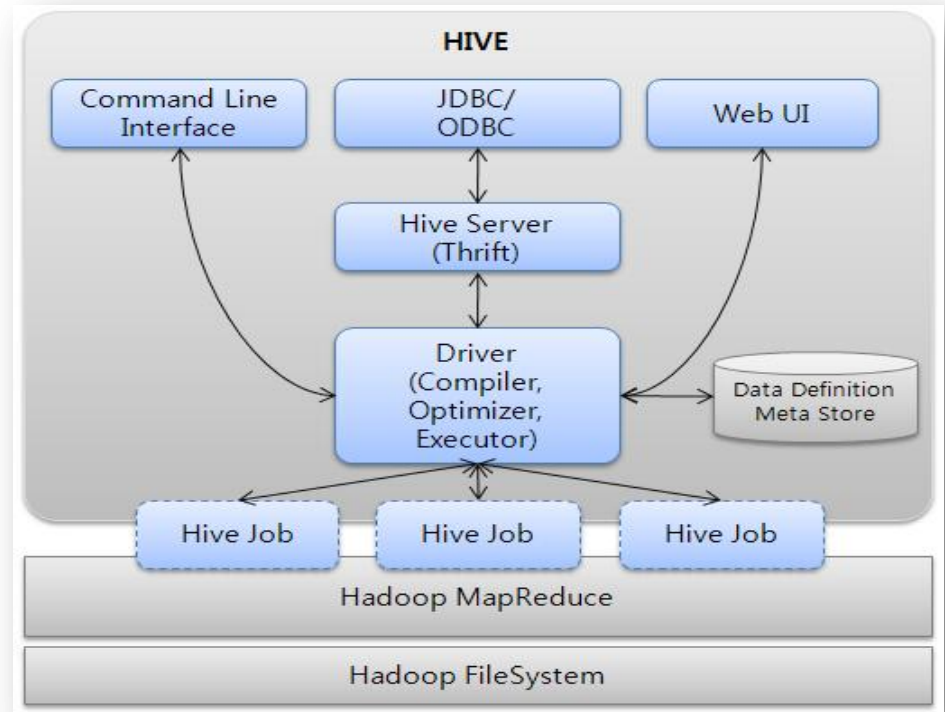
쉬운 데이터 접근 환경 구성

Crawler, ZooKeeper, Hadoop, Bamboo는 모두 개발자 관점
애초 필요했던 데이터 분석가가 쉽게
접근, 분석할 수 있는 기능은?



HIVE

HDFS에 저장된 텍스트 기반의
테이블을 데이터를 SQL을 이용하여
Map/Reduce 분산 병렬 작업을
수행하게 하는 플랫폼



인용 리트윗 추출 Query

```
insert overwrite table retweeted_key  
select transform(id, created_at, rt_id, text)  
using 'python extractRetweet.py'  
as (tweetkey, id, created_at) from default.twitter_hk;
```

```
Default
[~/work/program/hive/hive-0.8.1-bin]# bin/hive
WARNING: org.apache.hadoop.metrics.jvm.EventCounter is deprecated. Please use org.apache.hadoop.log.metrics.EventCounter in all t
he log4j.properties files.
Logging initialized using configuration in jar:file:/Users/babokim/work/program/hive/hive-0.8.1-bin/lib/hive-common-0.8.1.jar!/hi
ve-log4j.properties
Hive History file=/tmp/babokim/hive_job_log_babokim_201205150943_1023683445.txt
hive> select * from test limit 10;
OK
data11 10
data12 10
data13 40
data14 60
data15 10
data11 10
data12 10
data13 40
data14 60
data15 NULL
Time taken: 3.409 seconds
hive> explain select * from test limit 10;
OK
ABSTRACT SYNTAX TREE:
(C TOK_FROM (C TOK_TABREF (C TOK_TABNAME test))) (C TOK_INS
EXPR TOK_ALLCOLREF)) (C TOK_LIMIT 10)))
STAGE DEPENDENCIES:
Stage-0 is a root stage
STAGE PLANS:
Stage: Stage-0
Fetch Operator
limit: 10
Time taken: 0.107 seconds
hive>
```

Cloumon Enterprise

cloumon-enterprise/hive/query/main.do?query#

Cloumon hadoop HBASE ZooKeeper HIVE Cassandra 91.48/3566.25MB Setting LogOut

Hive Cluster

Add Delete

Connections

- Smileserv HDFS
 - bloter
 - default
 - election
 - feb
 - noil
 - profile
 - yonseil

Table

CreateTime	Table	Field Delimiter	Owner
2012-02-23 10:06:26	tweetkey	null	hadoop
2012-02-23 10:16:59	retweeted_key	null	hadoop
2012-02-23 23:27:16	profile1	null	hadoop
2012-02-25 22:36:24	election_tweets	null	hadoop
2012-02-26 09:33:17	not_election_tweets	null	hadoop
2012-03-05 20:10:52	tweets_feb	null	hadoop
2012-03-06 13:13:57	url_profile	null	hadoop

Column

Name	Type	Comment
tweet...	string	
id	string	
creat...	bigint	

Storage Info

Smileserv HDFS-1 Smileserv HDFS-2

Query Executor

select * from election.tweetkey limit 10

Clear Explain Execute

History

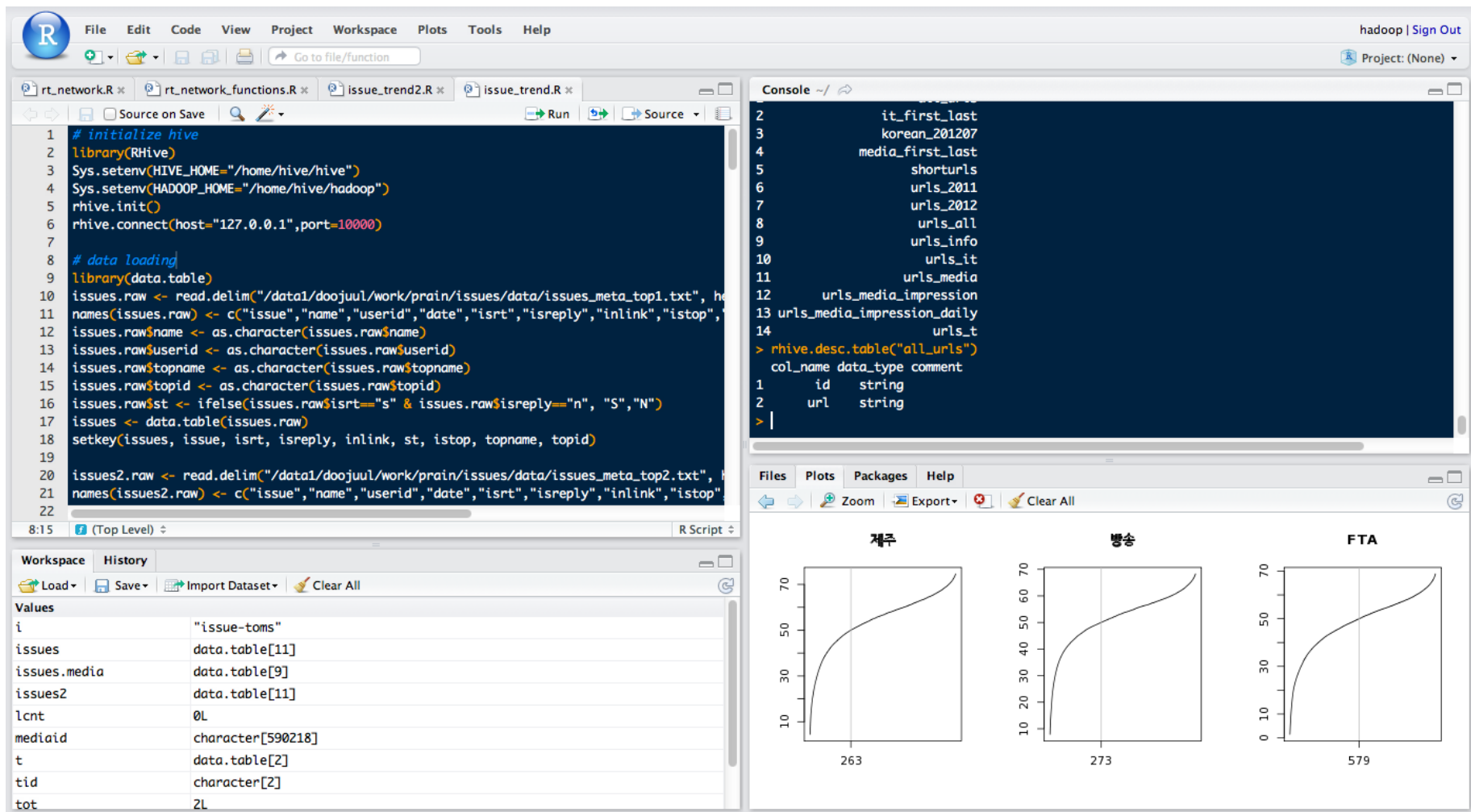
AutoLoad Reload

CreateTime	Tasks	Status
2012-04-14 01:55:50	0	success

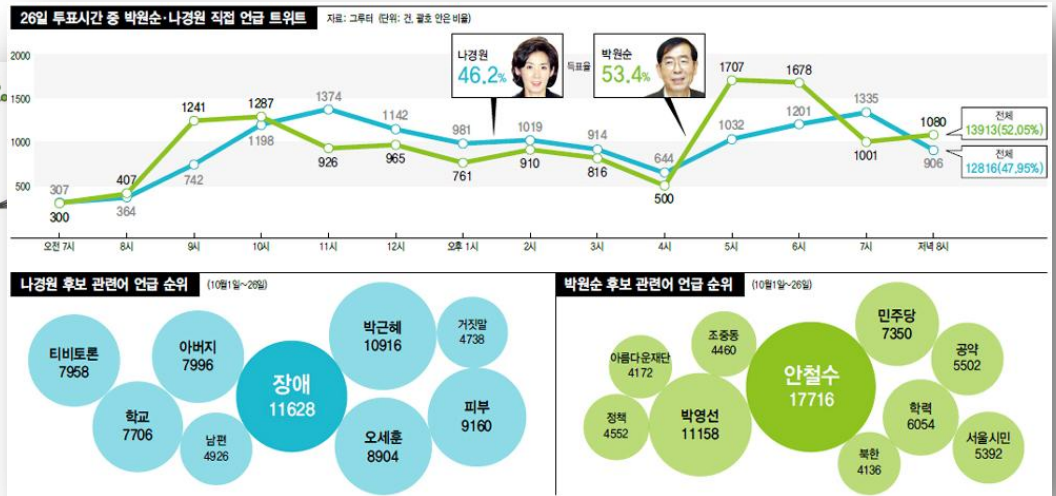
Result

Stop Reload

tweetkey	id	created_at
kurana_namhit...	1878730320507...	1333627396000
bemesmerized9...	1878730331370...	1333627396000
saintrecord_RT...	1878730349951...	1333627396000
69kunta69_@h...	1878730554716...	1333627401000
blockphoto_bo...	1878730650430...	1333627403000
woo_jung_0221...	1878730721231...	1333627405000
cr_990127_@B...	1878730829947...	1333627408000
dlrkf8803_@dl	1878730927213...	1333627410000

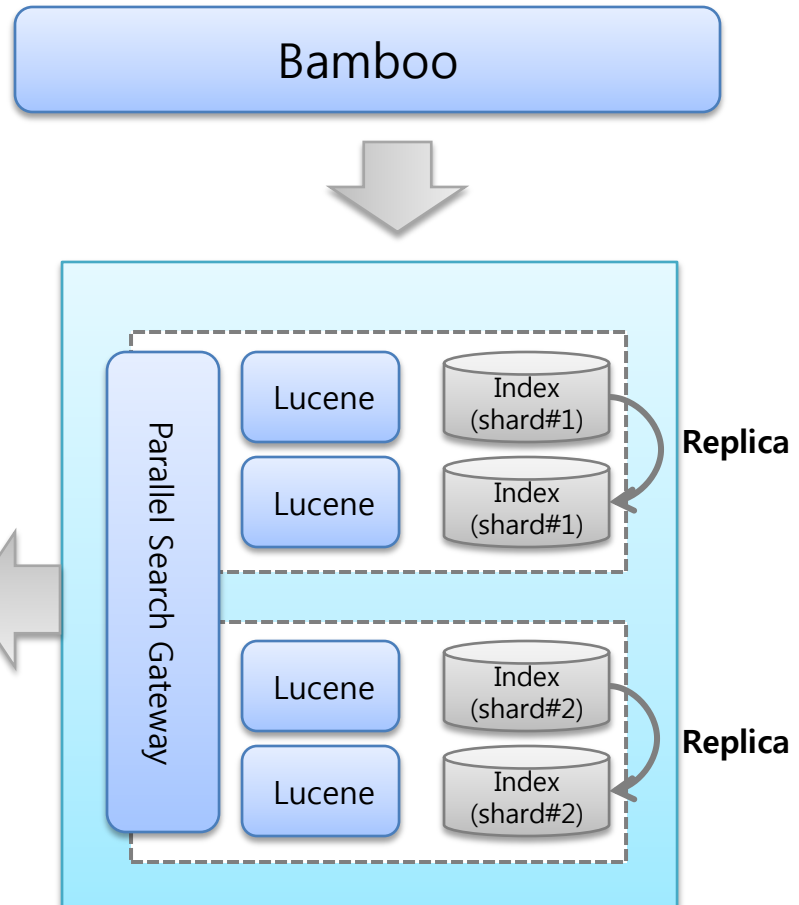


[블로터5th] ①1만명 중 4명, 하루도 안 쉬고 '재잘재잘' (블로터닷컴 2011년 9월 4일)
 [블로터5th] ③트위터 '1등 언론' 조사해보니 (블로터닷컴 2011년 9월 5일)
 '트위터 영향력' 100위중 박원순 지지자가 72명 (한겨레 2011년 10월 28일)
 소수 1%가 장악한 트위터...여소야대 여론 쏠림 극심 (한국경제 2011년 12월 26일)
 정치·기업·연예계 최고의 파워 트위터러는 이정희·삼성경제研·슈퍼주니어 (한국경제 2011년 12월 26일)
 보수·진보 따로 뭉치는 팔로어...트위터가 오히려 소통의 벽 (한국경제 2011년 12월 29일)
 뭉치는 진보 vs 탄죽거는 보수... 보혁 트위터 이용도 '땀땀' (한국일보 2012년 3월 22일)
 [총선] 공천 트위터 민심, 민주당에 더 부정적 (한국일보 2012년 3월 25일)
 민주당 떠나는 트위터 민심 (한국일보 2012년 3월 5일)
 [편집국에서/3월 6일] SNS서 '방자하다' 소리 듣는 민주당 (한국일보 2012년 3월 6일)
 '트위터 퍼나르기' 막판 젊은 표 몰렸다 (서울신문 2012년 4월 6일)
 野性の 중년男 'SNS총선' 이끈다 (서울신문 2012년 4월 6일)
 朴·文 트위터 핵심키워드 '안철수' (서울신문 2012년 4월 6일)
 민주, SNS 선거운동은 잘했는데... (한국일보 2012년 4월 12일)



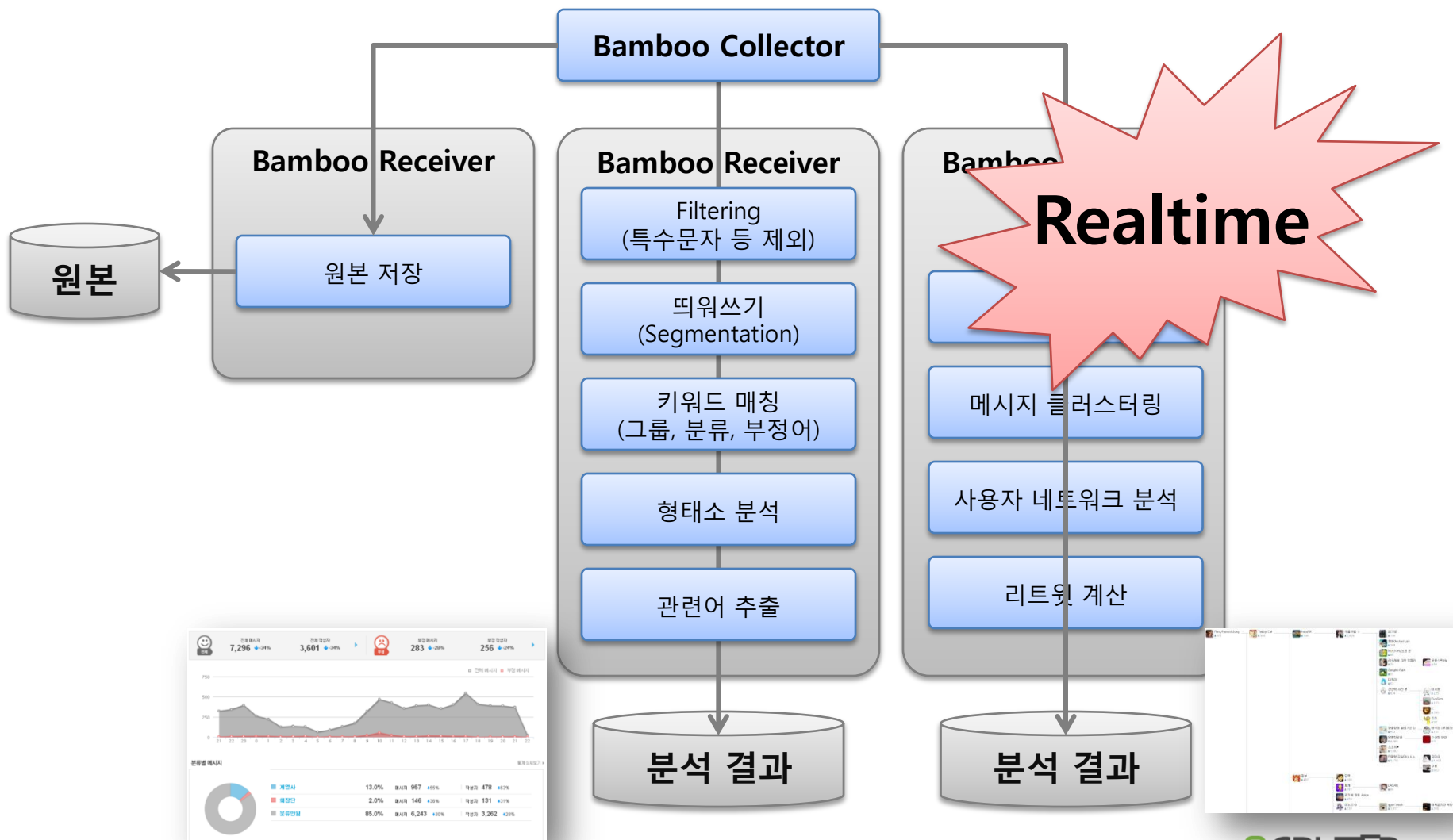
데이터를 구조적으로 저장해 보자... (검색엔진)

오픈 소스 Lucene 기반 분산 검색 구성
인덱스 볼륨 이중화 구성으로 장애 대응
크롤 즉시 검색 인덱스에 반영(실시간 검색)

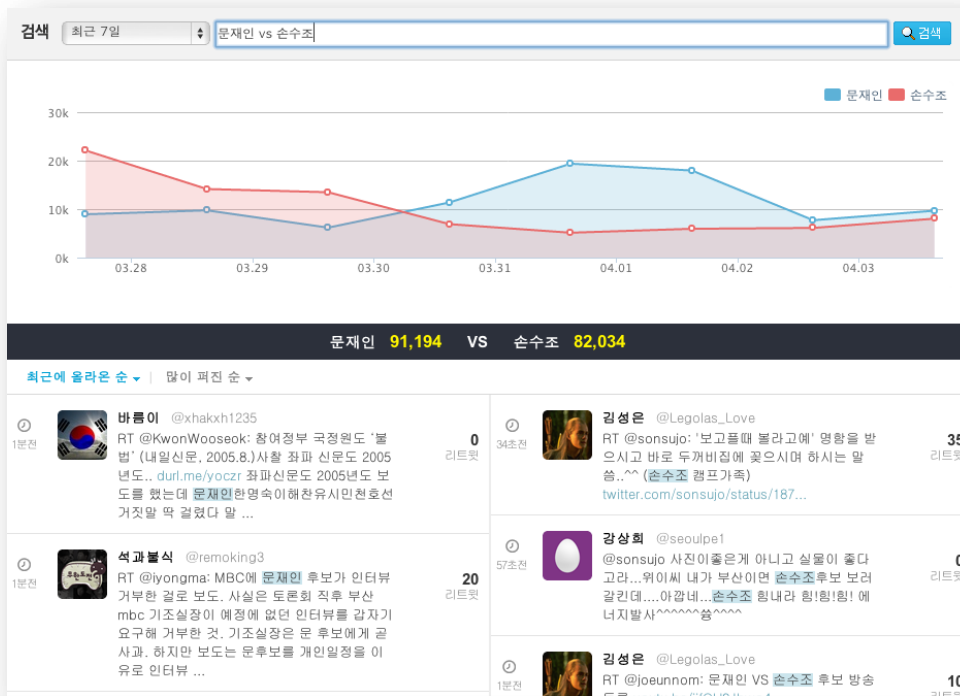


ElasticSearch

키워드 분석, 리트윗 경로 분석 등 실시간 분석
특정 사용자 메시지, 특정 키워드가 아닌 전체 메시지에 대한 분석



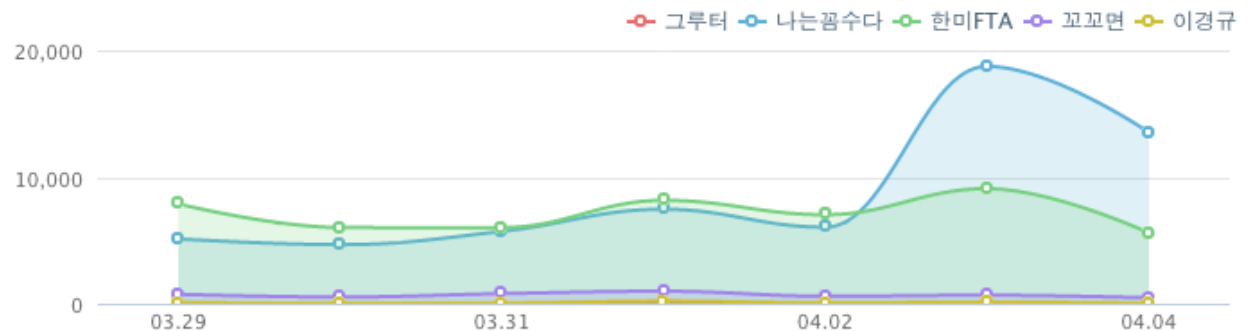
키워드 비교 검색



키워드 모니터링



메시지 트랜드



부정 메시지 트랜드



배치, 특정 트윗

6개월 소요

실시간, 모든 트윗

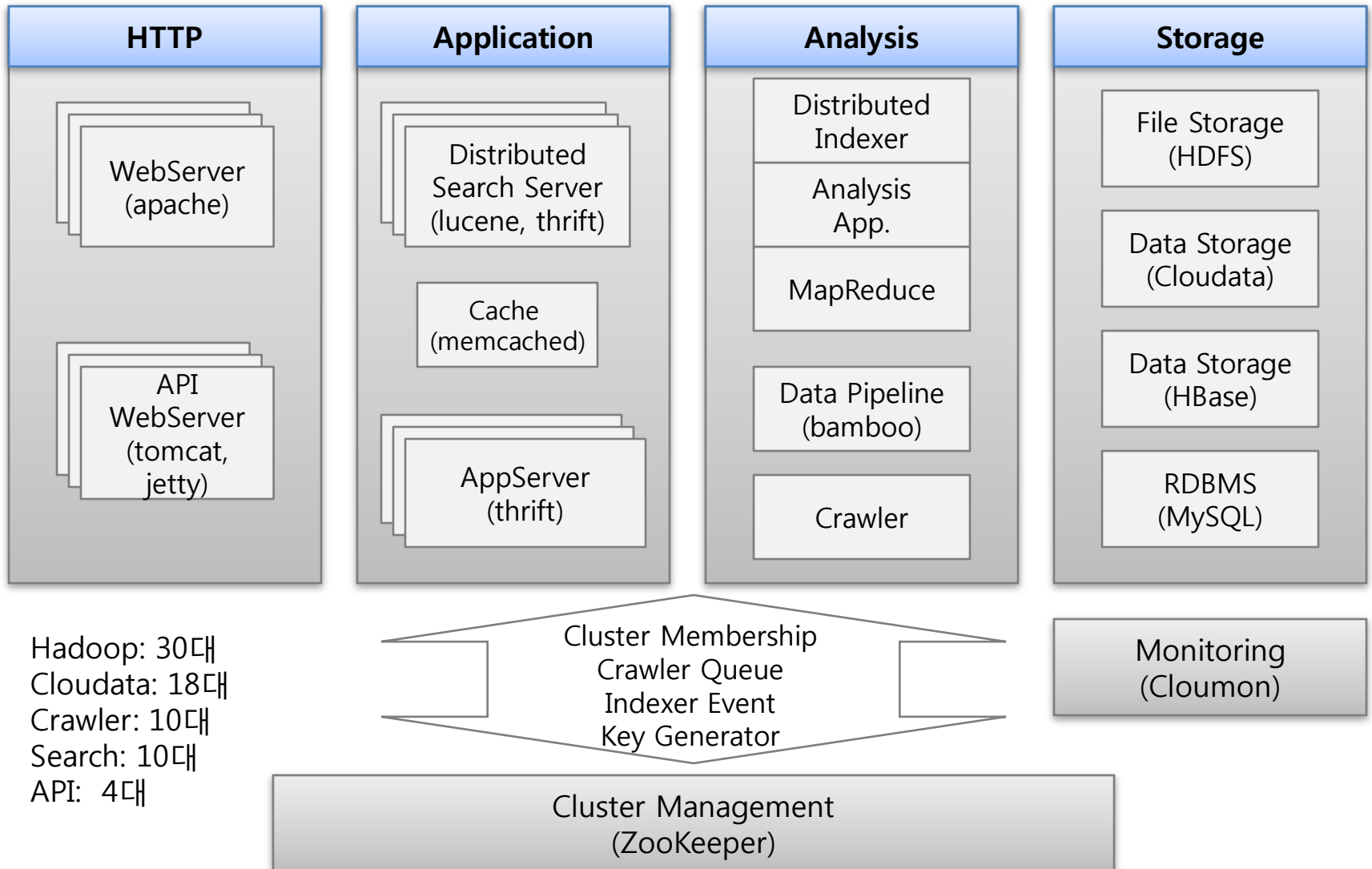


리트윗 정보				첫 트윗
5일전 첫트윗	899 리트윗	2일째 리트윗 지속	394,229 노출	 씨날은 빅데이터플래폼을 가지고 있는 회사가 새로운 서비스를 자연스럽게 만들어낼 수 있는 좋은 예입니다. "빅데이터 플랫폼 기반 소셜 분석 서비스 등장" http://www.ddaily.co.kr/news/news_view.php?uid=88496

YK, Kwon @gaiaville | 3월 7일, 18시 00분 첫 트윗



전체 시스템 구성



Hadoop: 30대
Cloudata: 18대
Crawler: 10대
Search: 10대
API: 4대

DEMO

감사합니다

contact@gruter.com

http://www.gruter.com

http://www.seenal.com