

Laboratório do Framework Hadoop em Plataformas de Cloud e Cluster Computing

Eng. André Luís Tibola.

Prof. Dr. Cláudio Fernando Resin Geyer.

Mst. Julio César Santos dos Anjos.

Junior Figueiredo Barros.

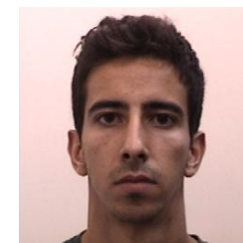
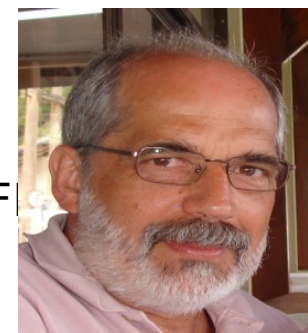
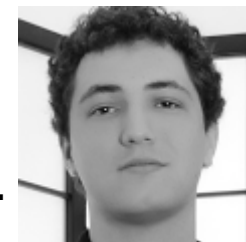
Mst. Raffael Bottoli Schemmer.

Apresentação dos autores

- GPPD – **G**ruppo de **P**rocessamento **P**aralelo e **D**istribuído.
 - <http://www.inf.ufrgs.br/gppd>.
- Principais linhas de pesquisa do GPPD/SLD:
 - **BigData** analytics.
 - Computação pervasiva e distribuída (**Ubicomp**).

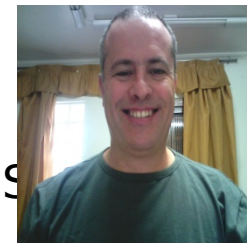
Apresentação dos autores

- Eng. André Luis Tibola
 - Mestrando em Ciência da Computação PPGC/UFRGS.
- Prof. Dr. Cláudio Fernando Resin Geyer
 - Professor Associado do Instituto de Informática – UFRGS.
- Junior Figueiredo Barros
 - Cursa engenharia da computação INF/UFRGS.
 - Pesquisador IC do GPPD.



Apresentação dos autores

- Prof. Mst. Julio César Santos dos Anjos
 - Doutorando em Ciência da Computação PPGC/UFRGS



- Mst. Raffael Bottoli Schemmer
 - Mestre em Ciência da Computação – PUC/RS.
 - Aluno especial PPGC/UFRGS.



Apresentação dos autores

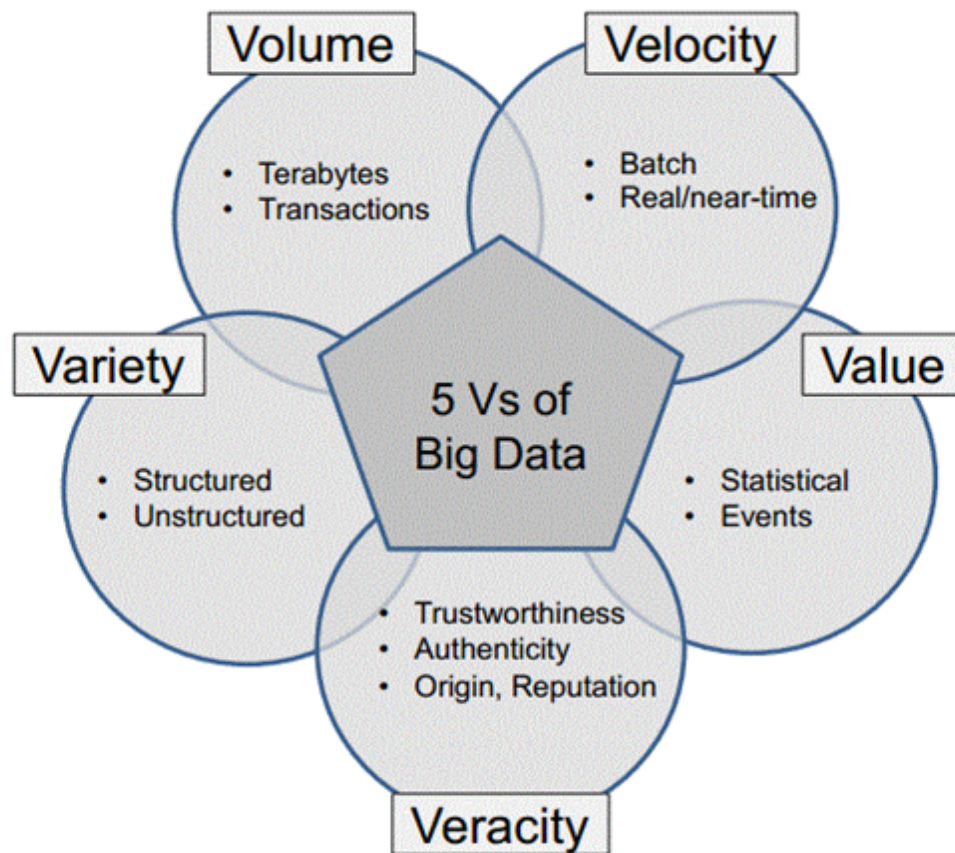


Apresentar o modelo de programação MapReduce e a utilização do framework Hadoop em ambientes de cluster e cloud computing.

Sumário

- Introdução ao BigData.
 - Estudo dirigido com ênfase ao Hadoop.
- Introdução ao Hadoop.
- Modelo de programação MapReduce.
- Arquitetura do Hadoop.
- Contextualização da infraestrutura:
 - O cluster GPPD GradeP.
 - A cloud Microsoft Azure.
- Instalação e configuração do Hadoop.
- Laboratório prático de programação.

Introdução ao BigData



Os 5Vs (Desafios) do BigData. "Marr (2015)".

Introdução ao BigData

Volume of Data
Generated by
Computers



Every 60 seconds



98,000+ tweets



695,000 status updates



11million instant messages



698,445 Google searches



168 million+ emails sent



1,820TB of data created

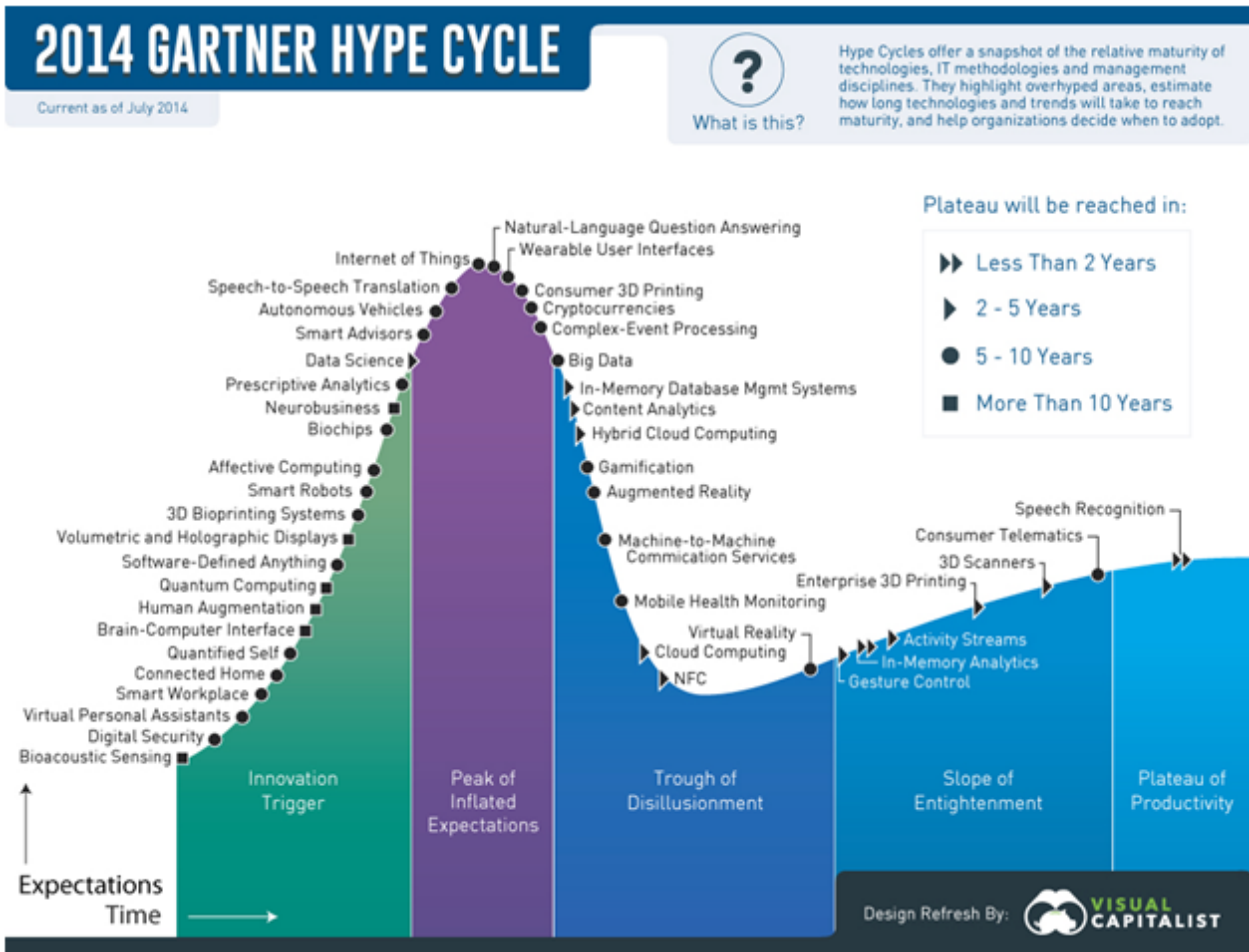


217 new mobile web users

Desafios de uma aplicação real. “Hp (2015)”.

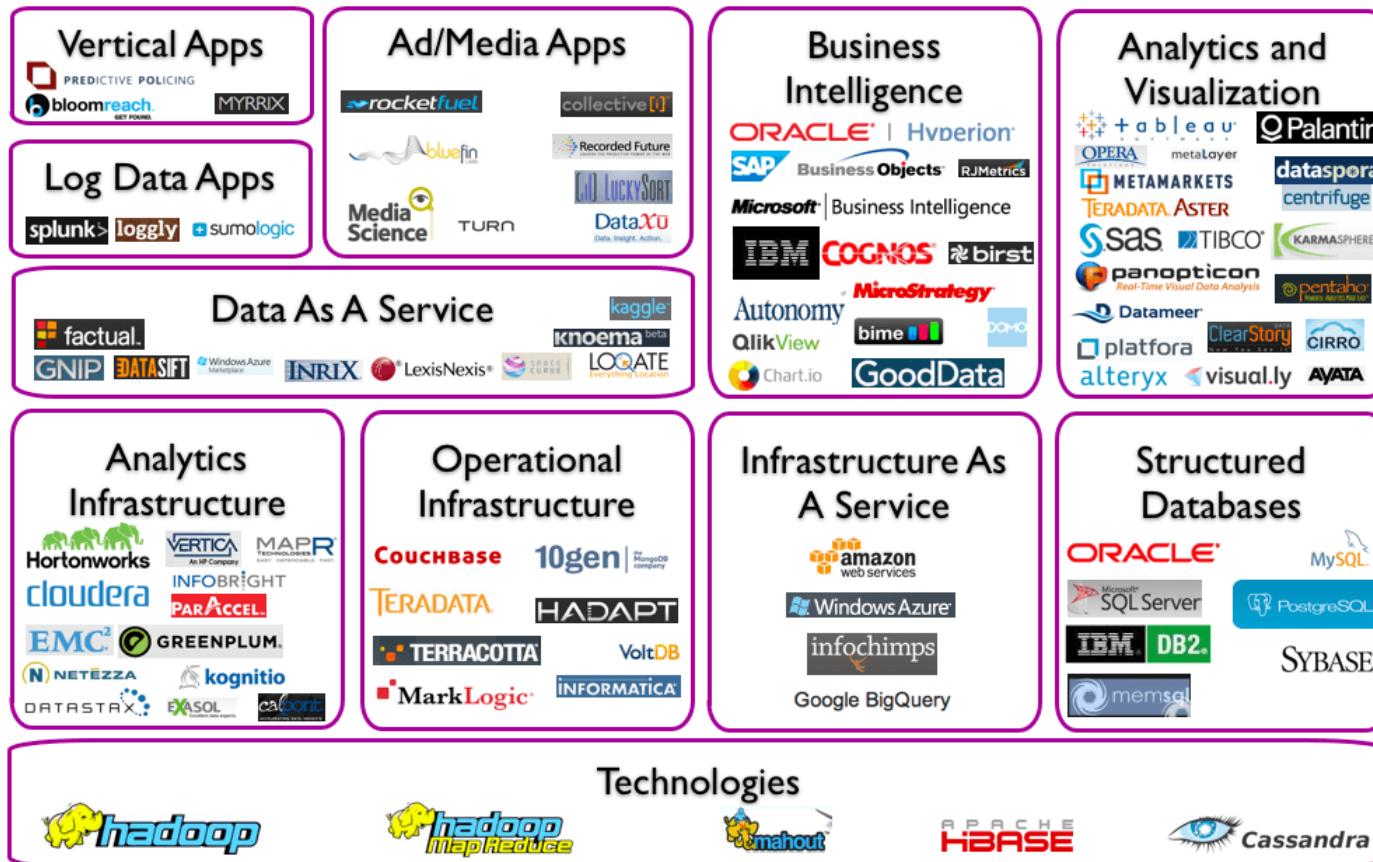


Introdução ao BigData



Visão do mercado sobre BigData na curva de *hype*. “Gartner (2014)”

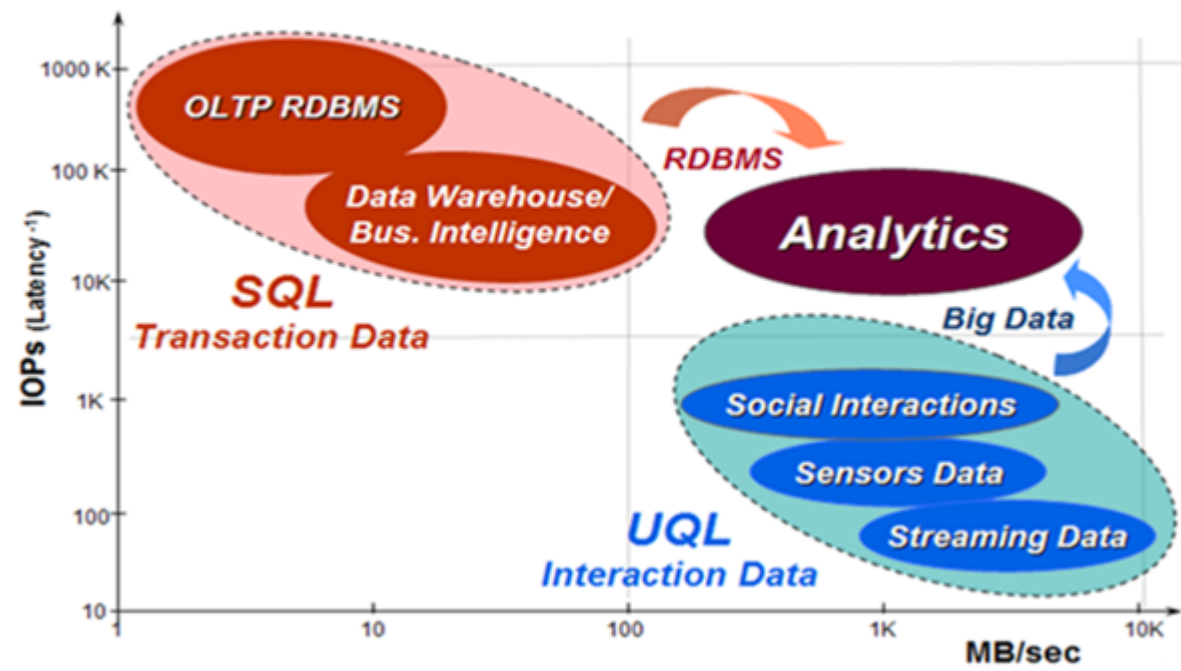
Introdução ao BigData



Soluções de BigData oferecidas pelo mercado. “Cognos (2013)”.

● Introdução ao BigData

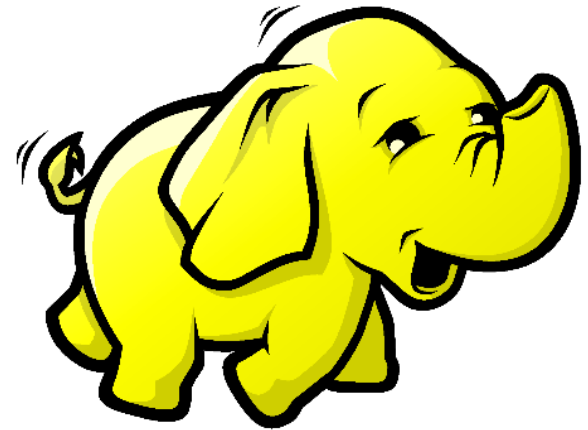
Analytics



Mundo real (SQL) Vs. Mundo desejado (UQL). “Imex (2014)”

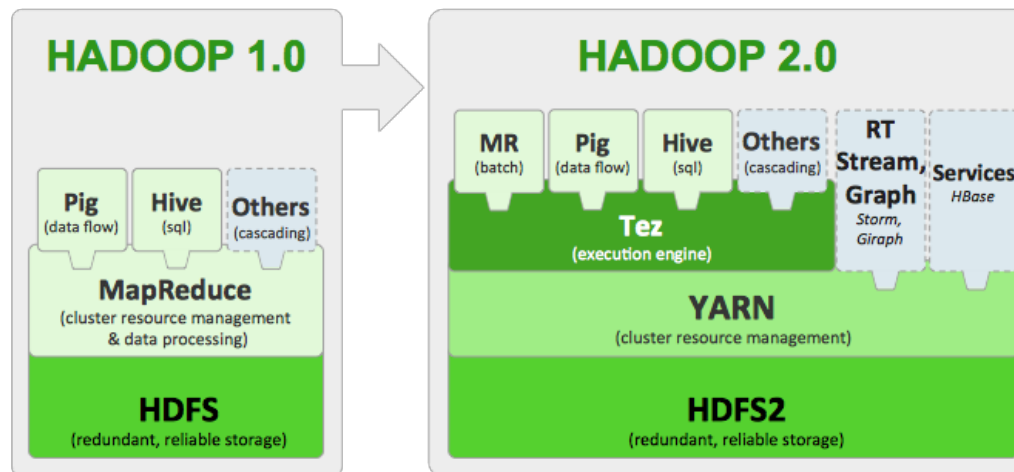
Introdução ao BigData

- Alternativa para o BigData:
 - Apache Hadoop Framework.
 - Ênfase deste minicurso.
- Razões para uso do Hadoop:
 - Uma das soluções mais aceitas/adota
 - Open Source (Não comercial).
 - Suporte a inúmeras APIs (Integração).
 - Multiplataforma (Java).
 - Framework modular (Extensível).
 - Suporte a várias linguagens de programação (Wrappers).
 - Escalabilidade de recursos:
 - Volume de dados Vs Adição de novos recursos.



Introdução ao Hadoop/MapReduce

- Criado no ano de 2005 (Nutch Project):
 - Modelo MapReduce (2004).
 - Sistema de arquivos HDFS (2007).
- Projeto Apache (2008) - <http://hadoop.apache.org>.
- Primeira release 1.0.0 (2011).
- Hadoop 2.0 (2013) - Última versão 2.7.0 (2015).



Arquitetura inicial 1.X Vs. Arquitetura atual 2.X

Introdução ao Hadoop/MapReduce

- Versões de terceiros (MapReduce Comercial):



- HortonWorks.
- MapR.
- Cloudera.
- Microsoft HDInsight (HortonWorks



- Empresas que utilizam (Hadoop Modificado):

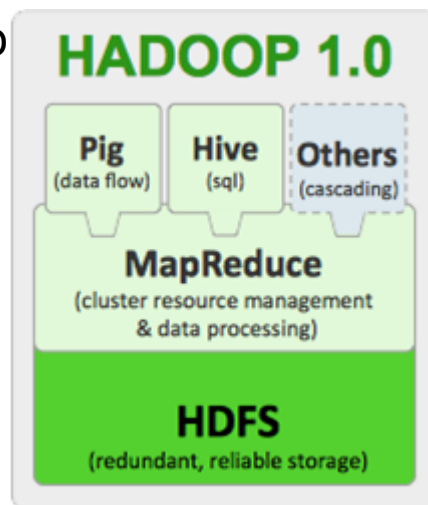


- Facebook.
- Yahoo.
- Amazon Elastic MapReduce (AWS - EMR).
- Pelo menos outras 50 empresas.



Introdução ao Hadoop

- Hadoop 1.2.1 (Hadoop 1.0) será trabalhado.
 - HDFS e MapReduce serão estudados em detalhes.
- Demais componentes:
 - Pig: Suporte a primitivas para dados não estruturados.
 - Hive: Suporte a primitivas para dados estruturados (HiveQL).
 - “Others” (Scoop) para queries SQL.



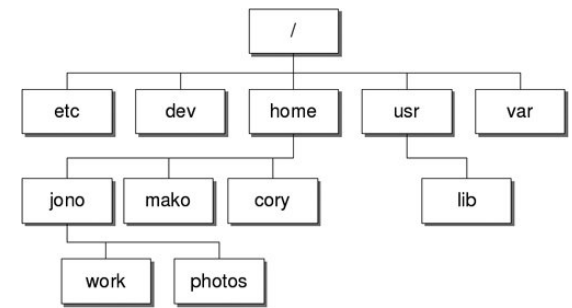
Arquitetura do Hadoop utilizada Apache (2013c).

Sistema de arquivos distribuído (HDFS)

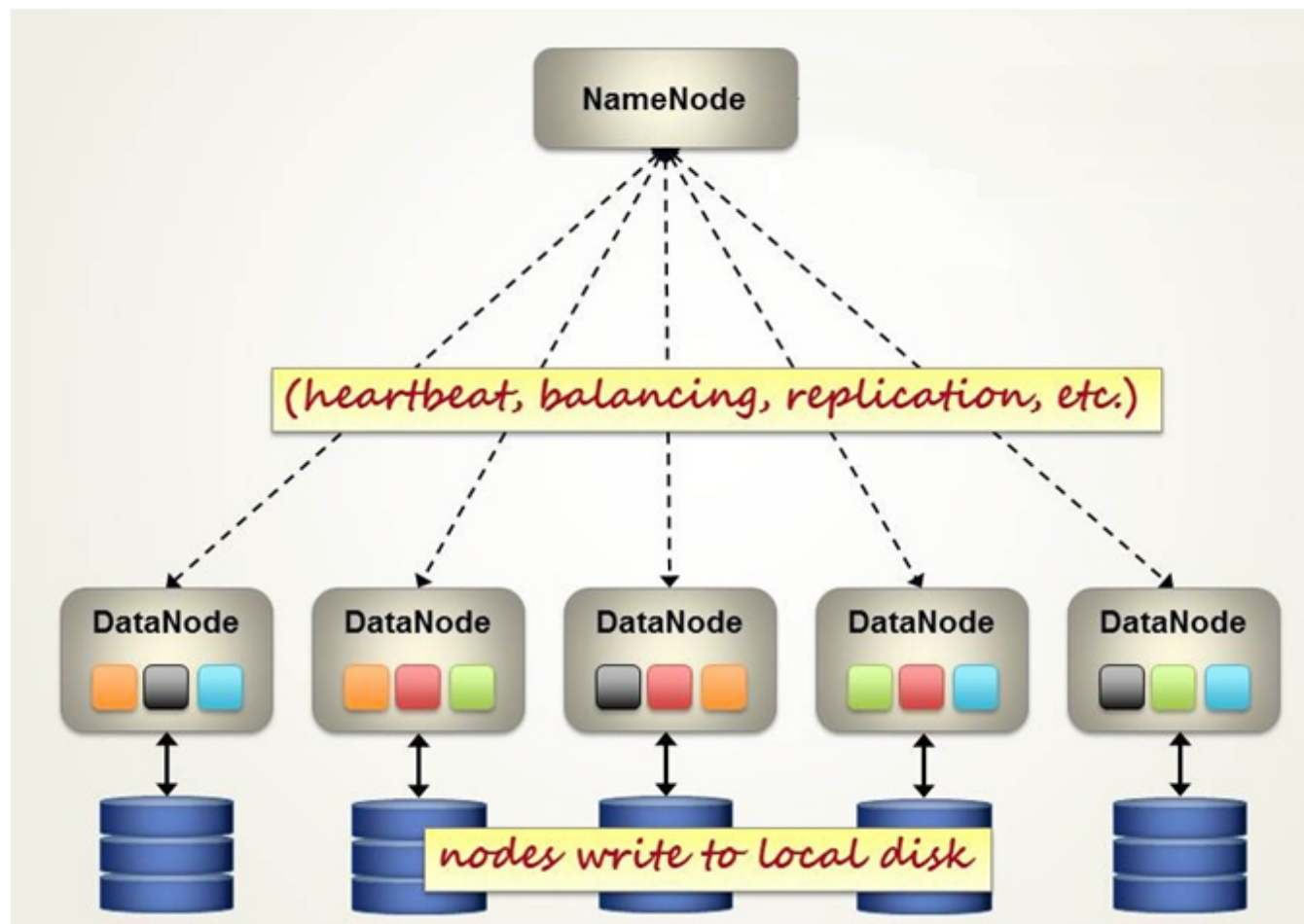
- **Hadoop Distributed File System (HDFS):**
 - Armazenamento de dados de forma distribuída.
 - Projetado para equipamentos de propósito geral.
 - Executa sob sistema de arquivo local.
- Tolerante a falhas:
 - Técnica implementada através da replicação de blocos.
 - Arquivos são fragmentados em pedaços (chunks).
- Implementação dirigida a leitura e escrita de altos volumes de informação (Largura de banda):
 - *Write Once Read Many.*

Sistema de arquivos distribuído (HDFS)

- O HDFS é implementado em uma arquitetura mestre/escravo:
 - (1) Namenode: Administra os dados.
 - (N) Datanode: Armazena os dados.
- Sistema de arquivos baseado em diretórios.
- Operações suportadas pelo HDFS:
 - Open, close, read, write.
 - Implementadas pelo Namenode.
- Namenode é responsável pela replicação dos dados:
 - Utiliza Heartbeat para controle.
 - Realiza balanceamento quando necessário.
- Datanode serve requisições.



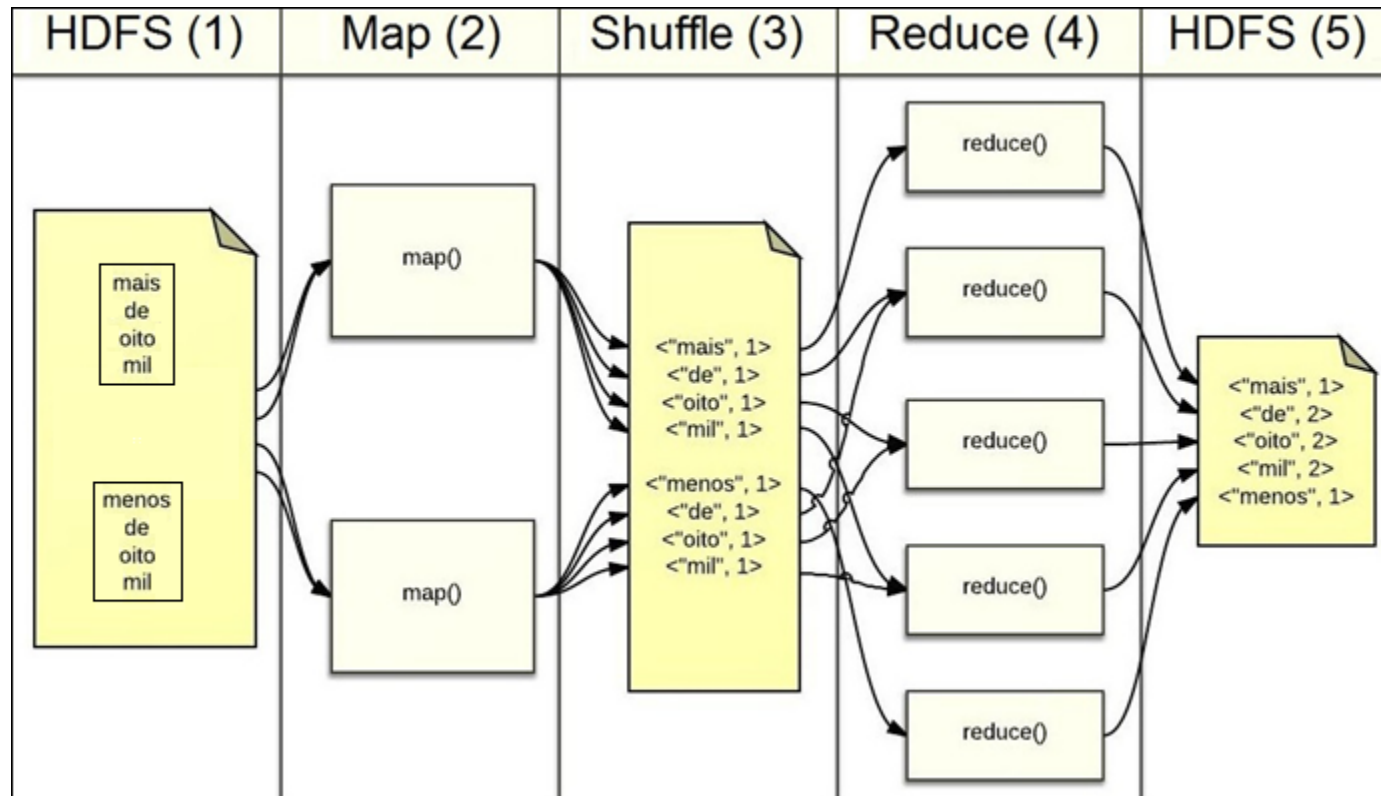
Sistema de arquivos distribuído (HDFS)



Modelo de programação (MR)

- Proposto em 2004 pela Google.
- Modelo de programação abstrato.
- Aplicado no processamento de dados distribuídos.
- Funcionamento do MapReduce:
 - Map: Processa todos os dados da entrada.
 - Reduce: Processa os resultados do Map.
 - Dados são transmitidos em tuplas <chave, valor>
- Processos MapReduce operam de forma independente.
- O framework trata questões inerentes de sistemas distribuídos como escalonamento e falhas.

Modelo de programação (MR)



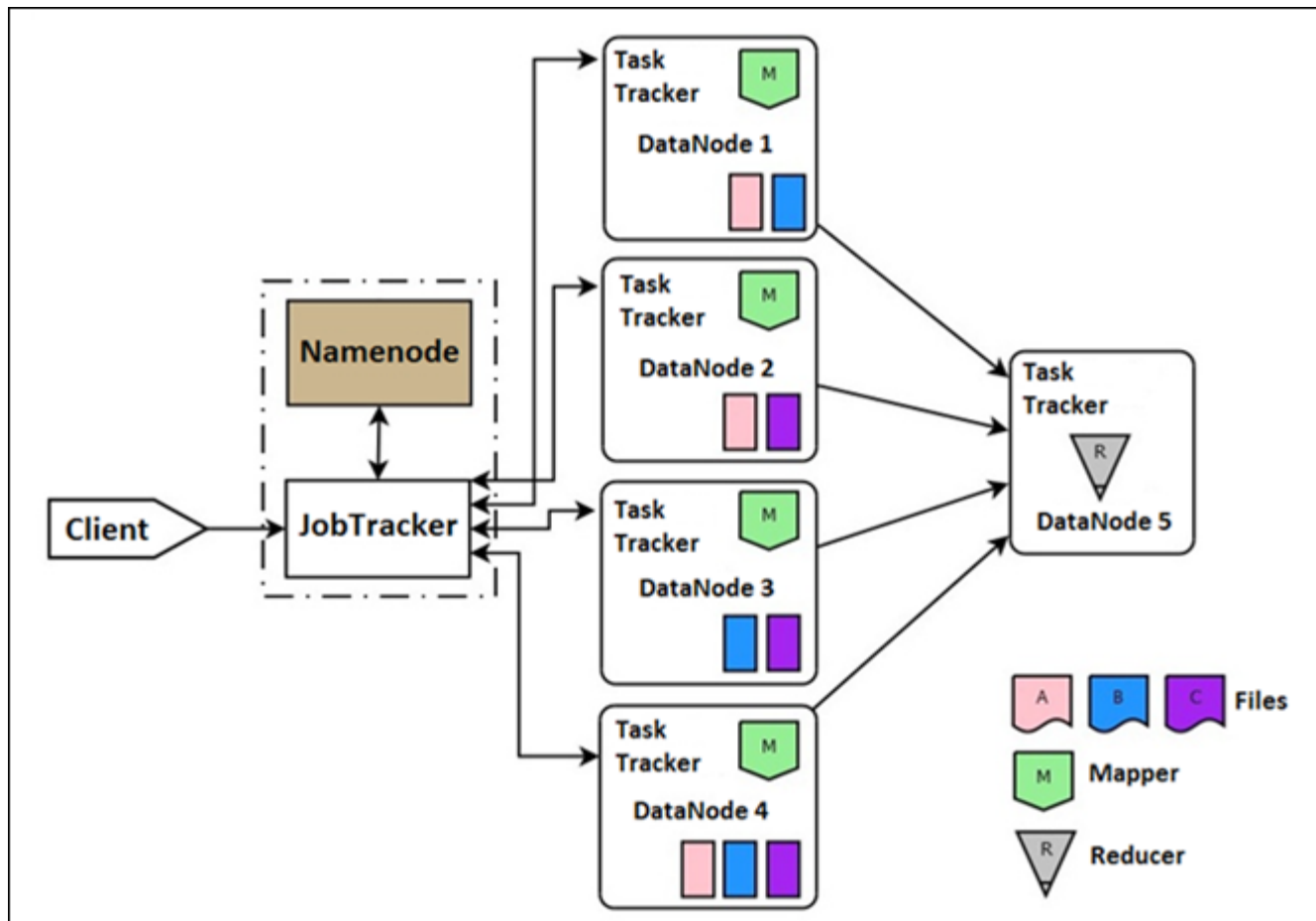
Modelo de programação MapReduce (WordCount)
Dean, J. and Ghemawat, S. (2004).



Modelo de programação (MR)

- O MR é implementado em uma arquitetura mestre/escravo:
 - (1) JobTracker: Administra os processos MR.
 - (N) TaskTracker: Executa as operações MR.
- MapReduce segue a abordagem de que:
 - A computação deve ser dirigida aos dados.
- Na visão do programador (usuário):
 - Aplicação deverá ser descrita como um código MapReduce.
 - HDFS deverá armazenar os dados da aplicação.
- Componente coordenador do MR (JobTracker):
 - Define o número de processos do tipo Map e Reduce:
 - Conforme o número de recursos.
 - Define o local onde os processos serão executados:
 - Conforme a localidade dos dados.

Modelo de programação (MR)



Arquitetura de componentes do Hadoop
Apache (2013a).



Recursos

- Execução especulativa
- Reescalonamento em caso de falhas nos workers
 - Map
 - Reduce
- Explora localidade dos dados
 - Local
 - Rack-Local
 - Remoto
- Pode evitar nós lentos



Resumo de conceitos

- Tecnologias como HDFS e MR são peças chave do Hadoop.
 - Demais APIs são construídas sob estas fundações.
 - Processos MR executam funções e rotinas destas APIs.
- HDFS framework responde pelas questões quanto:
 - Ao particionamento dos dados.
 - A saúde da informação (Replicação).
- MR framework responde pela questões quanto:
 - Ao mapeamento dos processos.
 - Ao escalonamento frente a mudança dos dados (Falhas).

Linhas de pesquisa em BigData do GPPD/SLD

- Aplicações de domínio específico.
- Capacity Planning.
- Garantia de nível de serviço.
- Green Computing.
- Infraestruturas heterogêneas.
- Infraestruturas híbridas.
- MR em ambientes voluntários.
- Simulação.
- Stream Processing.
- Tolerância a falhas.

Publicações Recentes

- MRA++: Scheduling and data placement on MapReduce for heterogeneous environments.
 - ACM Future Generation Computer Systems.
- MRSG : A MapReduce simulator over SimGrid.
 - Elsevier Parallel Computing.
- Genetic Mapping of Diseases through Big Data Techniques.
 - ICEIS 2015
- A Toolkit for Simulating MapReduce in Hybrid Infrastructures,
 - WS Big Data – SBAC-PAD - Oct 2014



MRSRG : A MapReduce simulator over SimGrid.

- Motivações para utilização de simulação:
 - Acesso escasso a plataformas de larga escala
 - Execução global é muito difícil de ser controlada
 - Facilidade para prototipar e avaliar alterações no modelo MapReduce
- MRSRG:
 - Desenvolvido pelo GPPD/SLD
 - Contempla os aspectos gerais do MapReduce
 - Capacidade de simulação de sistemas grandes



Dúvidas ?

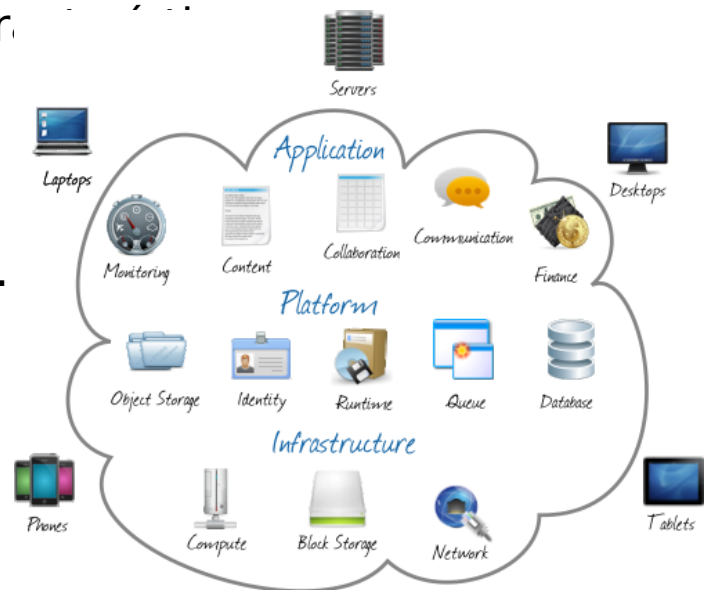
15:30
(Segunda parte)
Laboratório
prático e experimental
do Hadoop

Ambientes distribuídos

- Revisão conceitual:
 - Hadoop é projetado para em ambientes distribuídos.
 - Resiliente a falhas (Replicação de componentes).

- Ambientes distribuídos e suas características (atualidade):

- Clouds:
 - Máquinas compartilhadas.
 - Virtualização de recursos.
- Clusters:
 - Máquinas dedicadas.



O cluster GPPD GradeP

- GradeP: grade.p.inf.ufrgs.br.
- Cluster formado por computadores de propósito geral:
 - Recursos heterogêneos (tecnologias).
 - Rede ethernet de propósito geral.
- GradeP (201):
 - Execução do Hadoop em modo agendado.
 - 12 nós (computadores dedicados):
 - 1x Vostro 270s – Hadoop Master:
 - I5 3470s | 8GB RAM | 1TB HDD
 - 11x Optiplex GX270 – Hadoop Slaves:
 - P4 2.8GHz HT | 2.5GB RAM | 1TB HDD
 - Rede ethernet gigabit dedicada.



GPPD GradeP (2015).

Instalação do Hadoop (GPPD GradeP)

- Conjunto de etapas necessárias (9):
 - [1] Acessar nó principal (grade.inf.ufrgs.br).
 - Linux e Java já instalados.
 - Acesso SSH garantido entre todas as máquinas.
 - [2] Acessar nó Hadoop mestre (compute-0-0).
 - Download e extração Apache Hadoop 1.2.1.
 - [3] Configuração de arquivos (hadoop/conf).
 - [3.1] *core-site.xml* – Configurações do namenode.
 - [3.2] *hadoop-env.sh* – Java path.
 - [3.3] *hdfs-site.xml* – Número de réplicas do HDFS.
 - [3.4] *mapred-site.xml* - Configurações do jobtracker.
 - [3.5] Arquivos master e slaves.

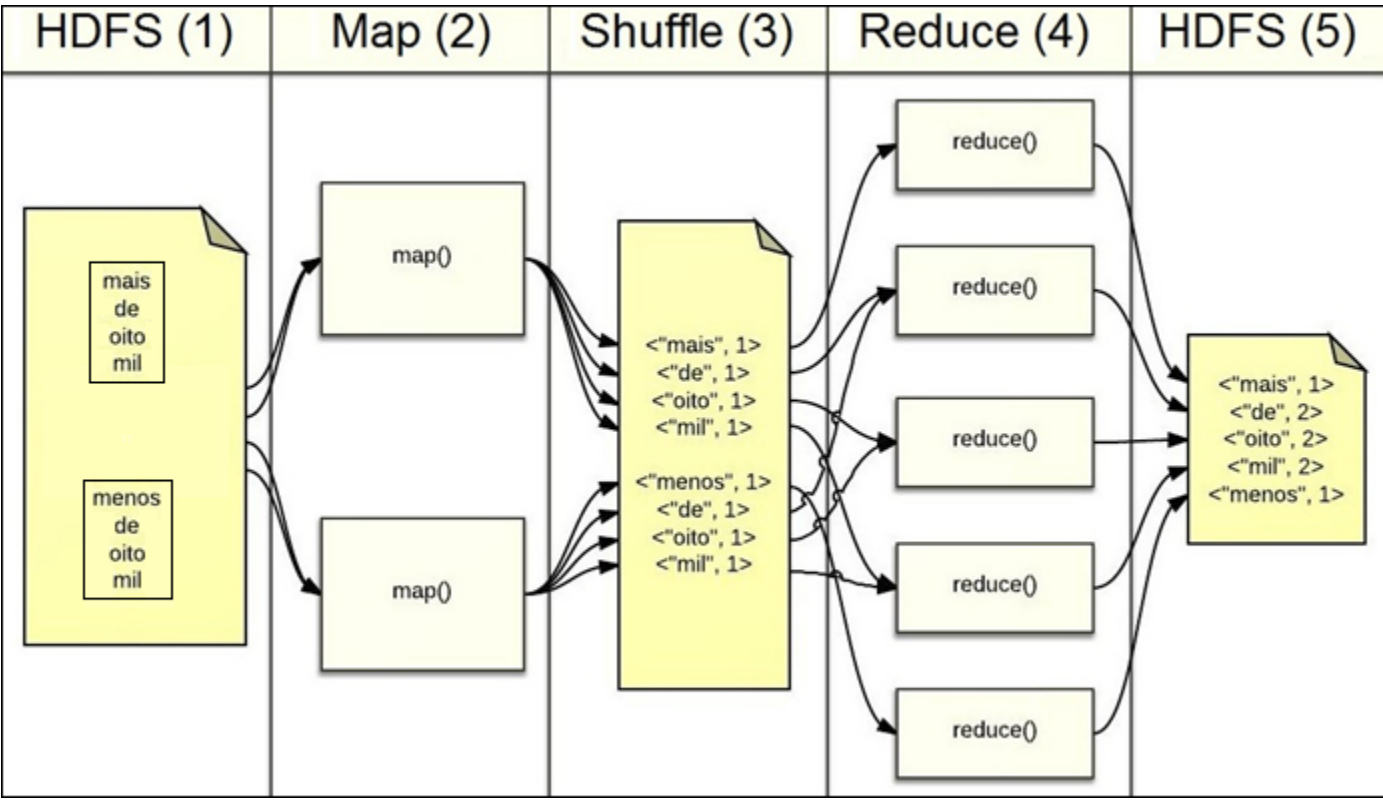
Instalação do Hadoop (GPPD GradeP)

- [4] Diretório HDFS e permissões.
- [5] Configuração bashrc (.bashrc).
- [6] Permissões de diretório do Hadoop.
- [7] Cópia do Hadoop.
- [8] Formatação inicial do HDFS.
- [9] Inicialização do Hadoop.

Laboratório Experimental (DFSAdmin)

- Dados deverão ser transferidos para o computador mestre (compute-0-0).
 - Este trabalho irá criar volumes de dados no mestre.
- Os arquivos deverão ser escritos no HDFS.
- O HDFS é acessível através do DFSAdmin.
 - API de acesso do usuário administrador do Hadoop.
- Principais comandos de manipulação do HDFS (DFSAdmin):
 - *hadoop dfs -ls*: Lista um diretório
 - *hadoop dfs -mkdir*: Cria um diretório
 - *hadoop dfs -copyFromLocal*: Escreve em um HDFS.
 - *hadoop dfs -copyToLocal*: Lê de um diretório no HDFS.
 - *hadoop dfs -cat*: Lê de um arquivo no HDFS.

● Contador de palavras (WordCount)



Funcionamento da aplicação de contagem de palavras em etapas (WordCount).



Geração de Trending Topics sob volumes de palavras.
"Timoe (2014)"

Laboratório Experimental (Word Count)

- Procedimentos a serem executados na GradeP:
 - [1] Estudo do código Java (WordCount)
 - [2] Escrita dos dados da aplicação no HDFS.
 - [3] Compilação e execução do código para 5 computadores.
 - [4] Análise dos resultados da execução.
 - [5] Demonstração da escalabilidade do Hadoop.
 - [6] Demonstração da resiliência (Tolerância a Falhas).



A cloud Microsoft Azure

- Microsoft Azure: azure.com
- Conjunto (finito) de recursos virtualizados:
 - Hyper-V (VMM): Gerencia a alocação dos recursos.
- Cada VM possui uma quantidade de recursos quanto:
 - Cores (1core a 32 cores).
 - RAM (1.75Gbytes a 448Gbytes)
 - Disco (10Gbytes a 6144Gbytes de SSD).
 - Rede (1Gbps Ethernet ou 10Gbps Infiniband).
- Azure trabalha com 3 níveis de recursos (serviços):
 - IaaS: Infraestrutura como um serviço (VM).
 - PaaS: Plataforma como um serviço (VM + SDK).
 - SaaS: Software como um serviço (VM + SDK + App).





A cloud Microsoft Azure

- 3 tipos de VMs (Máquinas/HW):
 - A (1-10): Propósito geral.
 - D (1-14): Alto desempenho.
 - G (1-5): Computação em dados.
- Este trabalho irá utilizar a Azure no nível de IaaS.
 - Com máquinas do tipo D1.
- HDDs devem ser definidos separadamente:
 - VMs podem suportar N instancias de discos.
 - Cada disco poderá ter até 1TByte de tamanho.
 - Discos são conectados diretamente as VMs (SANs).



Instalação do Hadoop (Microsoft Azure)

- Conjunto de etapas necessárias (10):
 - [1] Criação da cloud e das VMs + discos.
 - [2] Acessar nó principal (erad2015hadoop.cloudapp.net).
 - Instalação do Java nas VMs.
 - Instalação e acesso SSH entre as máquinas.
 - [3] Acessar nó Hadoop mestre (compute-0-0).
 - Download e extração Apache Hadoop 1.2.1.
 - [4] Configuração de arquivos (hadoop/conf).
 - [4.1] *core-site.xml* – Configurações do namenode.
 - [4.2] *hadoop-env.sh* – Java path.
 - [4.3] *hdfs-site.xml* – Número de réplicas do HDFS.

Instalação do Hadoop (Microsoft Azure)

- [5] Diretório HDFS e permissões.
- [6] Configuração bashrc (.bashrc).
- [7] Permissões de diretório do Hadoop.
- [8] Cópia do Hadoop configurado para os slaves.
- [9] Formatação inicial do HDFS.
- [10] Inicialização do Hadoop.

Laboratório Experimental (Word Count)

- Procedimentos a serem executados na Azure:
 - [1] Escrita dos dados da aplicação no HDFS.
 - [2] Execução do código para 5 computadores.
 - [3] Análise dos resultados da execução.
 - [4] Demonstração da escalabilidade do Hadoop.
 - [5] Demonstração da resiliência (Tolerância a Falhas).
 - [6] Comparação dos resultados GradeP Vs. Azure.



Fechamento do minicurso

- Vimos nesta aula que:
 - BigData: Conjunto de desafios da atualidade.
 - Mercado apresenta inúmeras soluções.
 - Hadoop é considerado como uma das soluções.
- Estudamos o Hadoop em nível conceitual (fundações):
 - HDFS e suas arquitetura de serviços.
 - O modelo de programação MapReduce.
 - Funcionamento do HDFS e do MapReduce.



Fechamento do minicurso

- Foi demonstrado na prática:
 - [1] Mecanismos básicos de operação do cluster e da cloud.
 - [2] Instalação e configuração do Hadoop.
 - [3] API DFSAdmin para administração do HDFS.
 - [4] Um código Java do tipo MapReduce (WordCount).
 - [5] Execução e visualização de resultados (WordCount).
 - [6] Escalabilidade do Hadoop.

Agradecimentos

- INF/UFRGS – GPPD:
 - Disponibilização do cluster GPPD GradeP.
 - Possibilitar a escrita e preparação do minicurso.
- XV ERAD 2015:
 - Viabilizar e fomentar o minicurso.
- Microsoft Research for Azure:
 - Disponibilização de acesso a Azure (IaaS).
- Público presente pela participação no minicurso.



Referências

- Bernard Mar. Big Data: Using Smart Big Data, Analytics and Metrics to Make Better Decisions and Improve Performance.
- Hp (2015) *"A BIG brother for your BIG data environment"*. Acessado em Abril de 2015. Disponível em: <http://h30499.www3.hp.com/t5/Business-Service-Management-BAC/A-BIG-brother-for-your-BIG-data-environment/ba-p/6284087>.
- Gartner (2014). Gartner Hyper Cycle. Disponível em : <http://www.gartner.com/technology/research/hype-cycles/>
- Cognos (2013). Big Data: Revolution or Hype?. Disponível em: <http://www.cognossource.com/big-data-revolution-or-hype/>
- Imex (2014). Big Data Industry Report 2014. Disponível em: <http://www.imexresearch.com/newsletters/BigDataEcosystem.html>



Referências

- Apache (2013b). MapReduce Tutorial - (WordCount v1.0). Apache.
- Apache (2013c). Overview of Apache Hadoop. Apache.
- Dean, J. and Ghemawat, S. (2004). MapReduce: Simplified Data Processing on Large Clusters. Commun. ACM, 51(1):107–113.
- Filho, B. (2013). Aplicação do MapReduce na Análise de Mutações Genéticas de Pacientes.



Referências

- GPPD (2015). Especificação e Documentação da GradeP. Grupo de Processamento Paralelo e Distribuído - GPPD/UFRGS.
- Ion-Life (2015). Ion Personal Genome Machine (PGM) System. Life Technologies.
- White, T. (2012). Hadoop - The Definitive Guide, volume 1. O'Reilly Media, Inc., 3rd edition.
- Timoe (2014). Word Cloud of Big Data. Disponível em:
<http://timoelliott.com/blog/wp-content/uploads/2014/01/big-data-speech-bubble.jpg>

Perguntas ?

raffael.schemmer@inf.ufrgs.br

INF/GPPD - UFRGS

