

REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE
MINISTERE DE L'ENSEIGNEMENT SUPERIEUR ET DE LA
RECHERCHE SCIENTIFIQUE

UNIVERSITE BADJI MOKHTAR ANNABA
FACULTE DES SCIENCES DE L'INGENIORAT
DEPARTEMENT INFORMATIQUE

HADOOP MAPREDUCE

Master : Gestion et Analyse des Données Massives (GADM)

2^{me} année

Dr. Klai Sihem

AVANT PROPOS. . .

Le Big Data est une science récente qui a surgie avec l'évolution et la variation des données et d'applications mises et échangées en ligne. Cette science consiste à prendre en charge d'une manière efficace un volume important des données hétérogènes en intégrant des techniques et outils nouveaux, vu que la technologie disponible ne répond plus aux besoins.

Ce polycopié est un support pédagogique qui permet d'initier l'étudiant au domaine des Big Data. Ce cours composé de plusieurs chapitres permet aux étudiants de comprendre la problématique et la motivation du domaine, et de maîtriser l'outil Hadoop avec le modèle MapReduce associé à ce domaine.

Chaque chapitre est élaboré pour répondre à un but pédagogique bien précis, se matérialisant par des explications, définitions accompagnées d'exemples et des illustrations par des figures suivies par des exercices, des solutions envisageables ou des fiches de travaux pratiques bien guidés.

1. Le chapitre I met l'étudiant dans le contexte du Big Data, consiste à lui donner des connaissances générales sur le domaine ;
2. Le chapitre II est consacré à l'étude de Hadoop, le framework qui permet le développement d'applications traitant les données massives. Ce chapitre donne les notions les plus générales avec la procédure d'installation du logiciel ;
3. Le chapitre III détaille la partie qui s'occupe du stockage des données "HDFS", avec la possibilité de la manipulation de ces données selon deux manières différentes à savoir : les commandes et l'API JAVA ;

4. Le chapitre IV étudie en détail la partie traitement des données massives "MapReduce", le modèle qui permet de traiter des blocs de données séparément et parallèlement dans des machines connectées. La modélisation selon le paradigme MapReduce est une étape importante avant le développement des programmes ;
5. Le chapitre V détaille l'implémentation des programmes MapReduce dans Hadoop. Dans le cadre de ce chapitre, nous étudions l'implémentation des programmes en utilisant le langage Java. D'autres langages peuvent être utilisés pour écrire des programmes mapreduce, mais cette partie n'est pas traitée dans ce cours.

L'élaboration de ce polycopié a été inspirée de plusieurs documents, j'ai pris le soin de les citer dans la partie bibliographie, d'autres documents aussi ont été cités afin d'apporter aux lecteurs plus de commandes et plus de détails sur les parties traitées dans ce cours.

TERMINOLOGIE. . .

JVM : Java virtuelle machine

HDFS : Hadoop Distributed File System

YARN : Yet Another Ressource Negotiator

AM : Application Master

API java : Application Programming Interfaces java

TABLE DES MATIÈRES

PRÉFACE	3
TABLE DES MATIÈRES	6
LISTE DES FIGURES	7
1 HADOOP	1
1.1 INTRODUCTION	2
1.2 HISTORIQUE DE HADOOP	2
1.3 PRÉSENTATION DE HADOOP	3
1.4 LES PRINCIPAUX COMPOSANTS D'HADOOP	4
1.5 LES MODES D'EXÉCUTION DE HADOOP	6
1.6 TÉLÉCHARGEMENT ET INSTALLATION DE HADOOP	6
1.7 PROGRESSION DES VERSIONS DE HADOOP ET INTÉGRATION DE YARN . . .	8
1.8 HADOOP STREAMING	9
1.9 EXERCICE DE COMPRÉHENSION	10
A ANNEXES	13
BIBLIOGRAPHIE	15

LISTE DES FIGURES

1.1	Les entreprises qui utilisent déjà Hadoop	3
1.2	L'architecture générale de Hadoop 03	4
1.3	Hadoop : architecture maître-esclave pour les données et les traitements 06	5
1.4	Architecture de Hadoop version 1. et 2 06.	9

HADOOP

Objectif Pédagogique

A la fin de ce chapitre, l'étudiant aura maîtrisé les concepts de base du logiciel Hadoop tels que son historique, ses principaux composants, ses modes d'exécution et son architecture. L'étudiant sera capable d'installer Hadoop dans un environnement Windows. Il pourra aussi faire la différence entre l'architecture maître-esclave classique et celle de Hadoop.

Pré-requis

L'étudiant doit avoir des connaissances générales sur l'architecture maître-esclave et le système distribué classiques ainsi que les connaissances acquises lors du chapitre ?? sur le Big data.

1.1 INTRODUCTION

HADOOP permet le stockage, l'analyse et le traitement des Big Data. C'est un logiciel open source de la fondation Apache et écrit en java.

Le framework Hadoop représente l'implémentation la plus populaire et la plus mature de l'approche conceptuelle MapReduce.

C'est un ensemble de classes écrites en Java pour la programmation des tâches MapReduce et HDFS.

Ces classes permettent au développeur d'écrire des programmes MapReduce sans que ce dernier ait à savoir comment les données et les traitements sont distribuées et parallélisées sur un ensemble de machines connectées.

Cet ensemble de machines connectées est appelé **cluster**.

Hadoop prend en charge les enjeux du traitement distribué tels que :

- L'optimisation des transferts disques et réseau en limitant les déplacements de données **data locality** ;
- La scalabilité pour permettre d'adapter la puissance au besoin **scalability** ;
- Et enfin la tolérance aux pannes **embracing failure** ^{o6}.

1.2 HISTORIQUE DE HADOOP

Hadoop est le nom du petit éléphant en peluche qui appartenait au fils de Doug Cutting, l'un des développeurs du framework Hadoop ¹. C'est pourquoi le logo de Hadoop est un éléphant.

- en 2003/2004 : Deux papiers publiés par Google, le premier sur un système de fichier distribué et le second sur le paradigme MapReduce pour le calcul distribué ;
- en 2004 : Développement de la première version du framework qui deviendra Hadoop par Doug Cutting et Mike Cafarella, ils s'inspirent de deux articles de recherches publiés par Google ;

1. <https://fr.slideshare.net/AmalAbid1/cours-big-data-chap2/>

- en 2006 : Ils développent une première version (au maximum quelques machines) exploitable de Apache Hadoop pour l'amélioration de l'indexation du moteur de recherche ;
- en 2008 : Hadoop est utilisé chez Yahoo dans plusieurs départements ;
- en 2011 : Hadoop utilisé par de nombreuses autres entreprises et universités. Le cluster Yahoo comporte 42000 machines et des centaines de peta-octets d'espace de stockage.

Plusieurs entreprises utilisent déjà Hadoop (voir Figure 1.1).



... et des milliers d'entreprises et universités à travers le monde.

FIGURE 1.1 – Les entreprises qui utilisent déjà Hadoop

1.3 PRÉSENTATION DE HADOOP

Hadoop contient plusieurs composants (voir figure 1.2) dont les principaux sont HDFS et MapReduce. Ils possèdent les caractéristiques suivantes :

1. Hadoop possède un système de fichiers distribué extensible. Il gère seul la distribution et le stockage des données sur les différents nœuds ou machines et pour augmenter la capacité de stockage, il suffit d'ajouter des nœuds dans la plateforme ;
2. Les nœuds de données sont aussi des nœuds de traitement, et par conséquent, l'augmentation des nœuds de données permet d'accroître la capacité de stockage ainsi que la capacité de traitement.

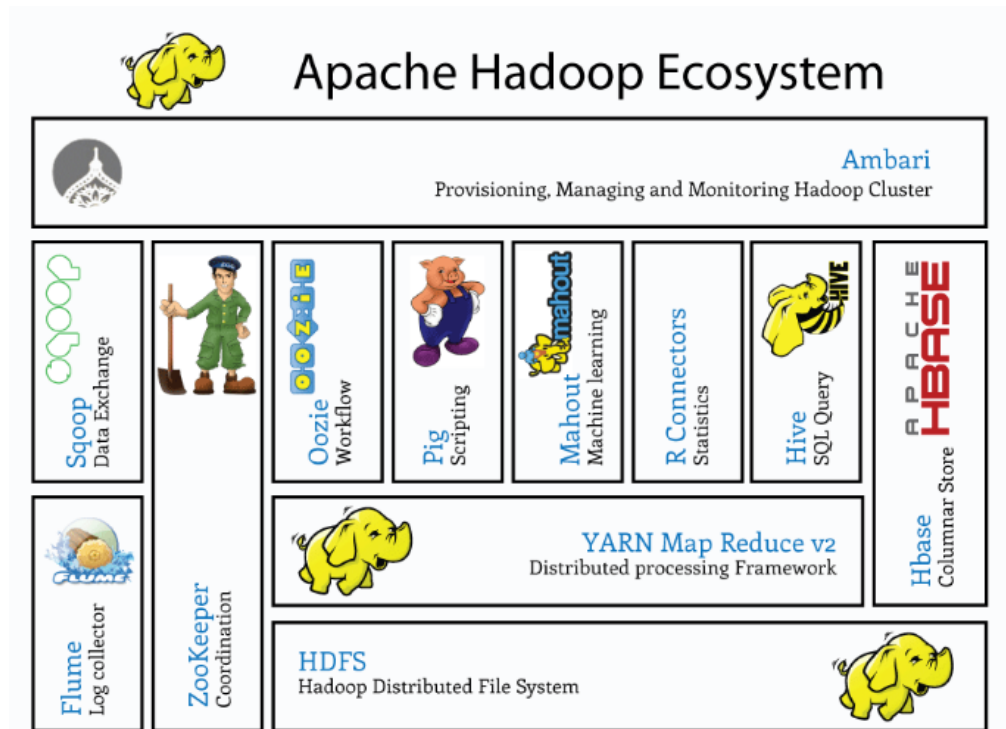


FIGURE 1.2 – L'architecture générale de Hadoop 03

3. Des mécanismes de tolérance aux pannes sont intégrés à la plateforme. Les données sont répliquées sur plusieurs nœuds afin d'être toujours accessibles. De même, les tâches de traitements exécutées sur les nœuds de données sont surveillées et relancées sur le nœud d'un autre répliquât si une panne survient.
4. Un paradigme de programmation MapReduce est intégré à la plateforme.

1.4 LES PRINCIPAUX COMPOSANTS D'HADOOP

Hadoop est composé essentiellement de deux parties importantes :

1. une partie stockage des données représenté par un système de fichiers HDFS, ce système se charge de la répartition des données sur plusieurs machines du cluster (voir les détails dans le chapitre ??), il est :
 - **Distribué** : les données sont réparties sur les machines du cluster ;
 - **Répliqué** : plusieurs copies des données sont stockées dans des nœuds différents ;

— Optimisé pour la colocalisation des données et des traitements.

Dans Hadoop, l'architecture de stockage est une architecture maître-esclave (voir figure 1.3) :

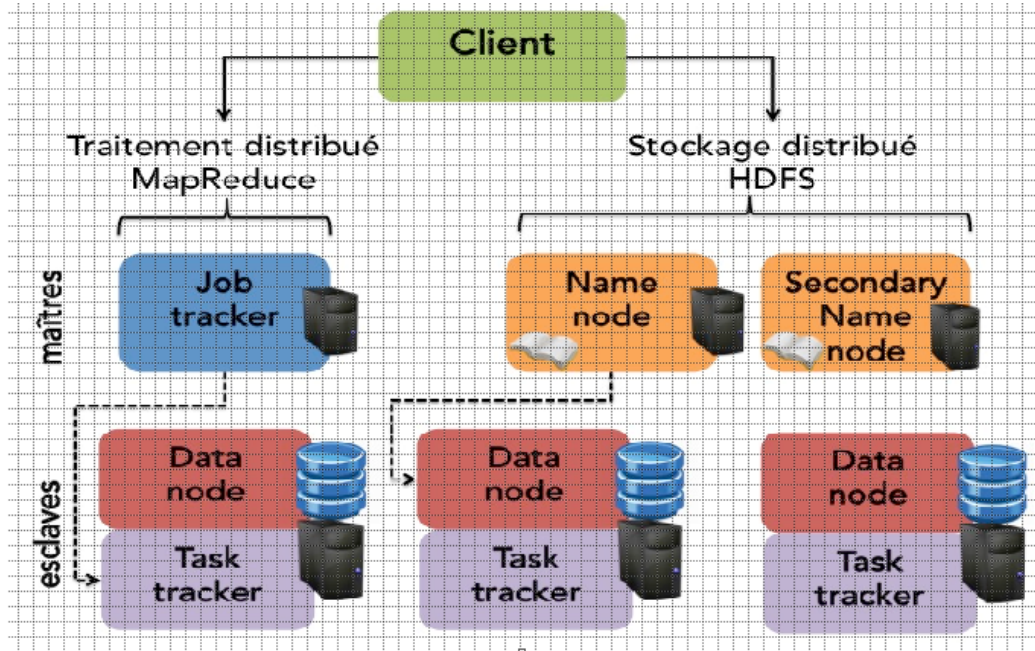


FIGURE 1.3 – Hadoop : architecture maître-esclave pour les données et les traitements 06

- Au niveau maître, le nœud appelé **name node** contient et stocke des métadonnées, c.a.d tous les noms des fichiers ainsi que leur localisation dans le cluster.
 - Le nœud appelé **secondary name node** sert de namenode de secours en cas de défaillance du nœud maître et il a donc pour rôle de faire des sauvegardes régulières des métadonnées ;
 - Au niveau esclave, les nœuds sont appelés **datanode** et contiennent les données, ils prennent en charge la gestion des opérations de stockage locales telles que la création, la suppression et la réplication des données sur instruction du namenode.
2. Une partie traitement représentée par MapReduce (voir figure 1.3 et chapitre ?? pour les détails), il permet :
- L'ordonnancement des traitements ;
 - La localisation des fichiers ;
 - La distribution de l'exécution.

Ces opérations sont transparentes au développeur, c'est le rôle de Hadoop avec une architecture de type maître-esclave composée de :

- Le **jobtracker** est un processus maître qui va se charger de l'ordonancement des traitements et de la gestion de l'ensemble des ressources du système.
- Le **tasktracker** est une unité de traitement du cluster. Il assure l'exécution et le suivi des tâches MAP ou REDUCE s'exécutant sur son nœud et qu'il reçoit du jobtracker.

1.5 LES MODES D'EXÉCUTION DE HADOOP

Hadoop propose trois modes d'exécution :

- **Mode local (standalone)** : dans ce mode, tout s'exécute au sein d'une seule JVM (Java Virtuel Machine). C'est le mode recommandé en phase de développement et de test des programmes.
- **Mode local pseudo-distribué** : dans ce mode, le fonctionnement en mode cluster est simulé par le lancement des tâches dans différentes JVM exécutées localement.
- **Mode distribué (fully-distributed)** : c'est le mode d'exécution réel d'Hadoop. Il permet de faire fonctionner le système de fichiers distribué et les tâches sur un ensemble de machines.

Remarque : Le mode local (standalone) est utilisé dans le reste de ce cours, pour les phases configuration et test des programmes (solutions des exercices).

1.6 TÉLÉCHARGEMENT ET INSTALLATION DE HADOOP

L'installation est particulière pour chaque système d'exploitation, Windows ou Unix.

Si on n'est pas sous Unix, on peut passer par une machine virtuelle² ou installer la plateforme Cygwin³ pour le cas de Windows.

Cygwin est un ensemble de packages de Unix portés sur Microsoft Windows. Il est nécessaire pour exécuter Hadoop sous Windows car il permet d'importer les packages de Unix dans Windows.

Cependant, l'installation d'Hadoop passe toujours par les mêmes étapes (les étudiants sont amenés à exécuter ces étapes lors des travaux pratiques :

1. L'installation et la vérification d'un ensemble de prérequis (c'est l'objet du TP N°1 en Annexes) :
 - Une version récente de Java doit être installée;
 - La création d'un groupe et d'un utilisateur spécifique à Hadoop (non obligatoire mais recommandé sous Unix);
 - Téléchargement et installation de la plateforme Cygwin pour le système d'exploitation Windows;
 - Configuration du service ssh pour permettre l'accès vers localhost pour l'utilisateur hadoop;
 - Démarrage de ssh localhost;
2. Téléchargement et installation d'Hadoop (c'est l'objet du TPN°2 en Annexes) :
 - Télécharger une version compressée de Hadoop à partir d'un site miroir⁴, par exemple : [hadoop-2.7.1.tar.gz](http://hadoop.apache.org/tarballs/hadoop-2.7.1.tar.gz);
 - Décompresser le fichier et installer Hadoop dans le répertoire de votre choix;
 - Se placer dans le répertoire d'installation d'Hadoop;
 - Indiquez l'emplacement de Java dans le fichier :
[etc/hadoop/hadoop-env.sh](#);

2. VirtuelBox. <https://www.virtualbox.org>

3. Cygwin. <http://www.cygwin.com>.

4. Hadoop. <https://www.apache.org/dyn/closer.cgi/hadoop/common/>

- Tester le fonctionnement de Hadoop en affichant la liste des paramètres acceptés par Hadoop.
- 3. Configuration de l'environnement Hadoop (TP N°2 en Annexes) : Pour la configuration de Hadoop, il faut modifier les fichiers de configuration suivantes :
 - Le fichier `.../hadoop/core-site.xml` permet de configurer le mode d'exécution de Hadoop ;
 - Le fichier `...etc/hadoop/hdfs-site.xml` contient les paramètres au système de fichiers HDFS ;
 - Le fichier `...etc/hadoop/mapred-site.xml` permet de configurer les paramètres spécifiques à MapReduce ;
 - Et enfin, le fichier `...etc/hadoop/yarn-site.xml` qui permet de paramétrer YARN.
- 4. Formater le système de fichiers HDFS local et démarrer Hadoop (voir TP N°2).

1.7 PROGRESSION DES VERSIONS DE HADOOP ET INTÉGRATION DE YARN

La première version de Hadoop présente des inconvénients, notamment pour les problèmes les plus complexes (qui nécessitent l'enchaînement de plusieurs séquences de MapReduce) à modéliser avec MapReduce et à exécuter sous Hadoop.

Cette complexité réside dans le fait que le jobtracker est trop chargé dans l'architecture de hadoop (version 1.x), il doit s'occuper de :

1. la gestion des ressources du cluster.
2. l'ordonnancement des jobs. La question qui se pose, et si le jobtracker est défaillant ?

Pour surmonter ces problèmes, des améliorations ont été apportées à Hadoop (version 2.x), où on a introduit YARN (Yet Another Resource Negotiator) dans l'architecture d'Hadoop :

- C'est un mécanisme qui permet de gérer des travaux (jobs) sur un cluster de machines ;
- Il permet aux utilisateurs de lancer des jobs Mapreduce sur des données présentes dans HDFS, et de suivre(monitor) leur avancement, récupérer les messages (logs) affichés par les programmes ;
- YARN peut déplacer un processus d'une machine à l'autre en cas de défaillance ou d'avancement jugé trop lent et il est transparent pour l'utilisateur ;
- YARN propose de séparer la gestion des ressources du cluster et la gestion des jobs Mapreduce.
- C'est un framework qui permet d'exécuter n'importe quel type d'application distribuée sur un cluster Hadoop, pas uniquement les applications Mapreduce, ceci en allouant les ressources des nœuds (mémoire et CPU) à d'autres applications si elles les demandent (voir la figure 1.4).

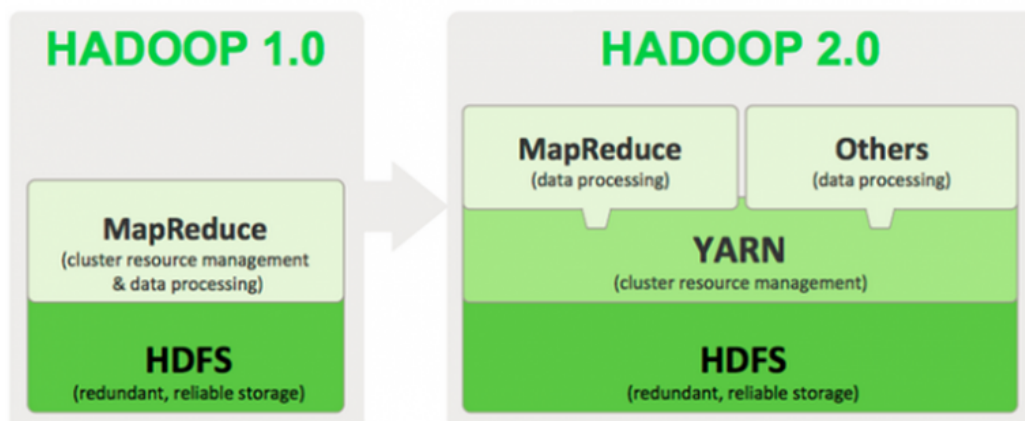


FIGURE 1.4 – Architecture de Hadoop version 1. et 2.06.

1.8 HADOOP STREAMING

Un autre concept important et utile pour le développement des programmes Mapreduce est Hadoop Streaming, c'est un outil fournis avec Hadoop, il permet l'exécution des programmes écrits dans d'autres langages que Java, tels que Python, C, C++.

1.9 EXERCICE DE COMPRÉHENSION

1. Donnez la bonne réponse : sont les composants principaux de Hadoop.
a : Hive
b : Hbase
c : HDFS
d : Flume
e : Mapreduce
2. Donner la bonne affirmation
a : HDFS s'exécute sur un petit groupe de nœuds.
b : Hadoop a besoin de matériels spécialisés pour traiter les données.
c : Le nombre de nœuds a de l'influence sur l'espace de stockage et la puissance du traitement.
3. Donner la bonne affirmation : La réplication peut intervenir dans quel cas :
a : Possibilité d'exécution des tâches sur d'autres nœuds. b : DataNode ne fonctionne plus.
c : Les blocs de données sont corrompus.
d : Le facteur de réplication est modifié.
e : Toutes les réponses sont bonnes.
4. Donnez la bonne réponse : peut être décrit comme un modèle de programmation utilisé pour développer des applications basées sur Hadoop.
a : HDFS
b : Mapreduce
c : Jobtracker
d : Cluster

5. Donnez la bonne réponse : est le nœud qui conserve les méta data.
- a : DataNode
 - b : NameNode
 - c : Cluster
6. Donnez la bonne réponse : Un nœud est responsable de l'exécution d'une tâche qui lui est assignée par le jobTracker :
- a : Mapper
 - b : taskTracker
 - c : jobTracker

ANNEXES

A

BIBLIOGRAPHIE

- [01]. site : Le big data, 2019. URL <https://www.lebigdata.fr/definition-big-data>.
- [02]. Stephane Vialle, CentraleSuplec. Big Data : Informatique pour les donnees et calculs massifs. 24 mai 2017 .
- [03]. mohammed zuhair al taie. hadoop ecosystem : an integrated environment for big data. in big data, guest posts 2015. URL <http://blog.agroknow.com/?p=3810>. (Cité page 4.)
- [04]. site : Hadoop. 2016. URL <http://www.opentuto.com/category/web-2/big-data/hadoop/>.
- [05]. amal abid. cours big data : Chapitre2 : Hadoop. 2017. URL <https://fr.slideshare.net/AmalAbid1/cours-big-data-chap2>.
- [06]. celine hudelot, regis behmo. realisez des calculs distribue sur des donnees massives. 2019. URL <https://openclassrooms.com/fr/courses/4297166-realisez-des-calculs-distribues-sur-des-donnees-massives>. (Cité pages 2, 5 et 9.)
- [07]. Benaouda, Sid Ahmed Amine, implantation du modele mapreduce dans l'environnement distribue, 2015 . URL <http://dspace.univ-tlemcen.dz>.
- [08]. Tom White, Hadoop : The Definitive Guide, June 2009 : First Edition.
- [09]. Srinath Perera, Thilina Gunarathne, Hadoop MapReduce Cookbook, First published : February 2013.

- [10]. Pierre Nerzic, Outils Hadoop pour le BigData, mars 2018.
- [11]. marty hall. map reduce 2.0 (input and output). 2013. URL <https://www.slideshare.net/martyhall/hadoop-tutorial-mapreduce-part-4-input-and-output>.
- [12]. build projects, learn skills, get hired. hadoop mapreduce- java-based processing framework for big data. URL <https://www.dezyre.com/hadoop-course/mapreduce>.
- [13]. URL <https://hadoop.apache.org/docs/r2.4.1/api/org/apache/hadoop/mapreduce/>.
- [14]. URL <https://gist.github.com/kzk/712029/9d0833aac03b23ec226e034d98f5871d9580724e>.
- [15]. Jeffrey Dean and Sanjay Ghemawat, MapReduce : Simplified Data Processing on Large Clusters, 2004 .
- [16]. Univ. Lille1- Licence info 3eme annee, Cours n°1 : Introduction à la programmation fonctionnelle, 2013-2014 .
- [17] dataflair team. hadoop architecture in detail hdfs,yarn et mapreduce. 2019. URL <https://data-flair.training/blogs/hadoop-architecture/>.
- [18] vlad korolev. hadoop on windows with eclipse. 2008. URL <http://v-lad.org/Tutorials/Hadoop/>.
- [19]. Donald Miner and Adam Shook. MapReduce Design Patterns. OReilly, 2012.
- [20]. Tom White. Hadoop : The Definitive Guide, 4th Edition. OReilly, 2015.
- [21]. Hadoop Training. 2018. URL <https://www.slideshare.net/AnandMHadoop/session-19-mapreduce>.