



BigData et traitement de données à l'échelle

MAXIME LEMAIRE – OLIVIER POTTIÉ – TIMELAB

Un peu d'histoire...

En 2006,  publie le whitepaper
“*Bigtable: A Distributed Storage System for Structured Data*”



voit le jour quelques années plus tard et confirme cette
nouvelle tendance

A large, billowing mushroom cloud from a nuclear explosion, with a bright white core at the base of the cap. The cloud is set against a dark, hazy background.

MEGADONNÉES !!



Avènement d'une nouvelle technologie :

Le NoSQL

pour

**Stocker de la donnée en masse
à la Google**

merci l'Open Source encore une fois quand même...



Bon en vrai ça ressemble à ça...

All in the NoSQL Family

NoSQL databases are geared toward managing large sets of varied and frequently updated data, often in distributed systems or the cloud. They avoid the rigid schemas associated with relational databases. But the architectures themselves vary and are separated into four primary classifications, although types are blending over time.



Document databases

Store data elements in document-like structures that encode information in formats such as JSON.



Common uses include content management and monitoring Web and mobile applications.



EXAMPLES:
Couchbase Server, CouchDB,
MarkLogic, MongoDB



Graph databases

Emphasize connections between data elements, storing related "nodes" in graphs to accelerate querying.



Common uses include recommendation engines and geospatial applications.



EXAMPLES:
Allegrograph, IBM
Graph, Neo4j



Key-value databases

Use a simple data model that pairs a unique key and its associated value in storing data elements.



Common uses include storing clickstream data and application logs.



EXAMPLES:
Aerospike, DynamoDB,
Redis, Riak



Wide column stores

Also called table-style databases—store data across tables that can have very large numbers of columns.



Common uses include Internet search and other large-scale Web applications.



EXAMPLES:
Accumulo, Cassandra, HBase,
Hypertable, SimpleDB

Particularités

- Différents types de datastore en fonction du besoin
- Pas de schéma ou de structure de données à priori
- Stockage distribué / répliqué
- Scalabilité horizontale
- Tolérance à la panne et haute disponibilité
- ACID vs BASE

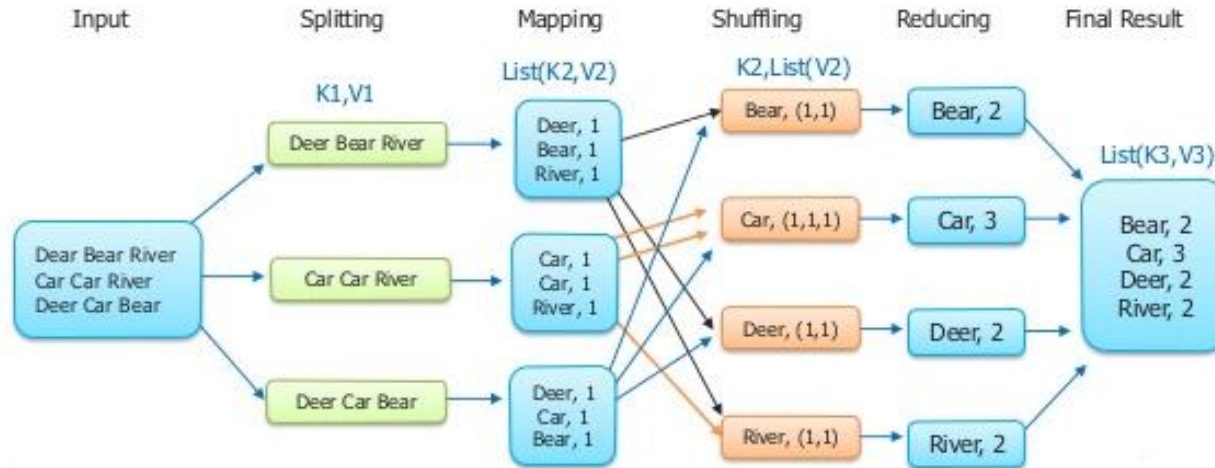
Atomicity, **C**onsistency, **I**solation, **D**urability VS **B**asically **A**vailable, **S**oft state, **E**ventual consistency



Et le traitement ?

MapReduce Paradigm

The Overall MapReduce Word Count Process



Et le traitement ?

Le MapReduce de Hadoop c'est bien mais :

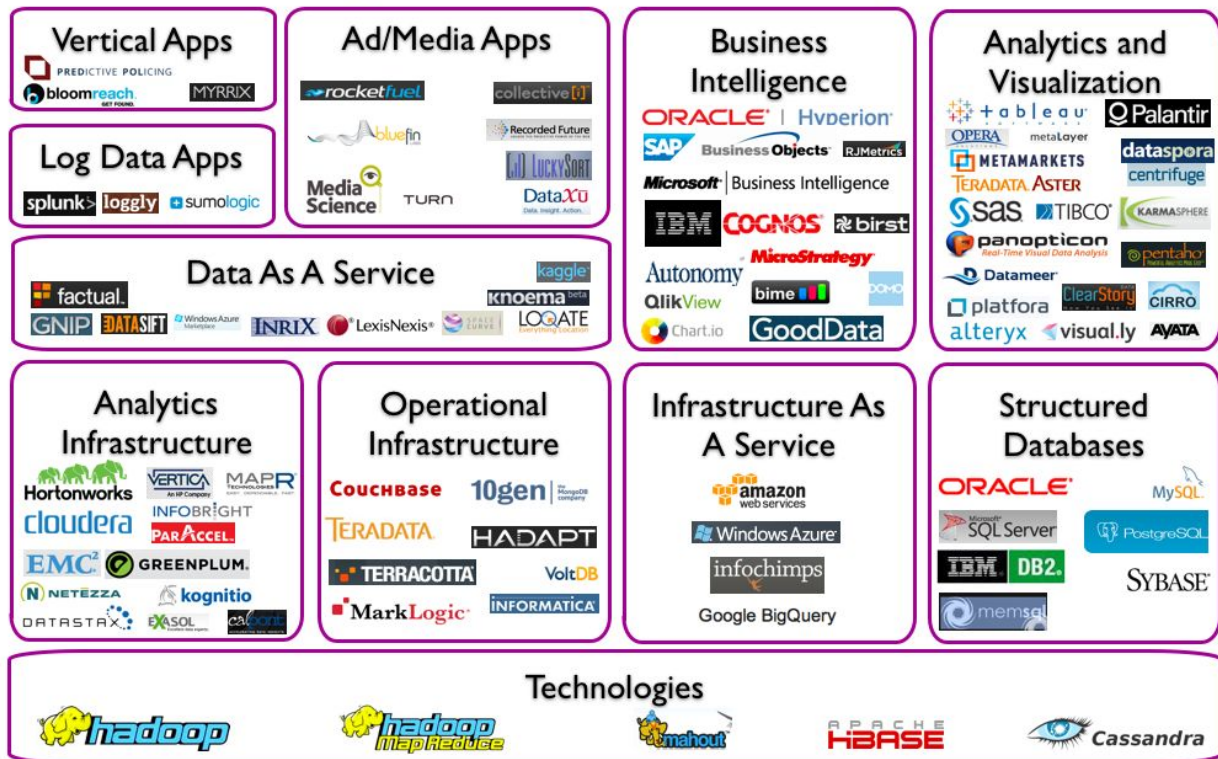
- c'est lent
- assez gourmand en ressources
- plutôt adapté au traitement batch et pas du tout au temps-réel

De nouveaux frameworks apparaissent, le plus notable



Le machine learning devient également une composante importante de ces nouveaux frameworks : il faut créer de la valeur !

BIG DATA LANDSCAPE 2012



En parallèle de cette évolution de la gestion de la donnée et de l'utilisation des ressources, la gestion de la couche applicative évolue également grâce aux containers et plus récemment aux orchestrateurs.

Chronologie de l'évolution de notre architecture



2013-2014

On teste toutes les technos qui nous semblent intéressantes :

Hadoop :(

MongoDB puis ElasticSearch pour les données utilisateurs

HBase, Cassandra, OpenTSDB, KairosDB pour les données de capteurs

Redis, Kafka (0.7 !) pour le messaging

On essuie les plâtres :

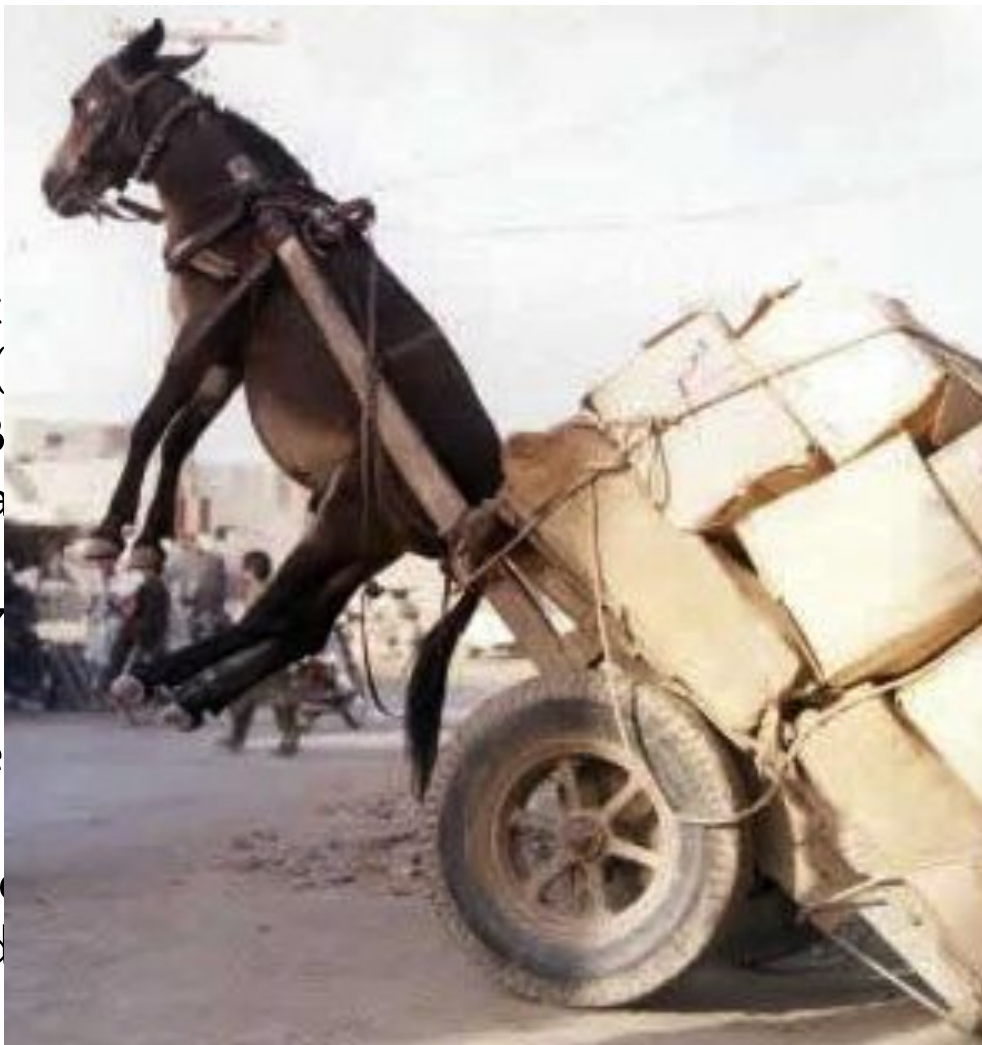
Les technos sont jeunes, pas toujours stables, complexes et on manque d'outil d'administration pour se simplifier la vie



On teste t
Hadoop :(
MongoDB
HBase, Ca
capteurs
Kafka (0.7

On essuie

Les techn
manque d

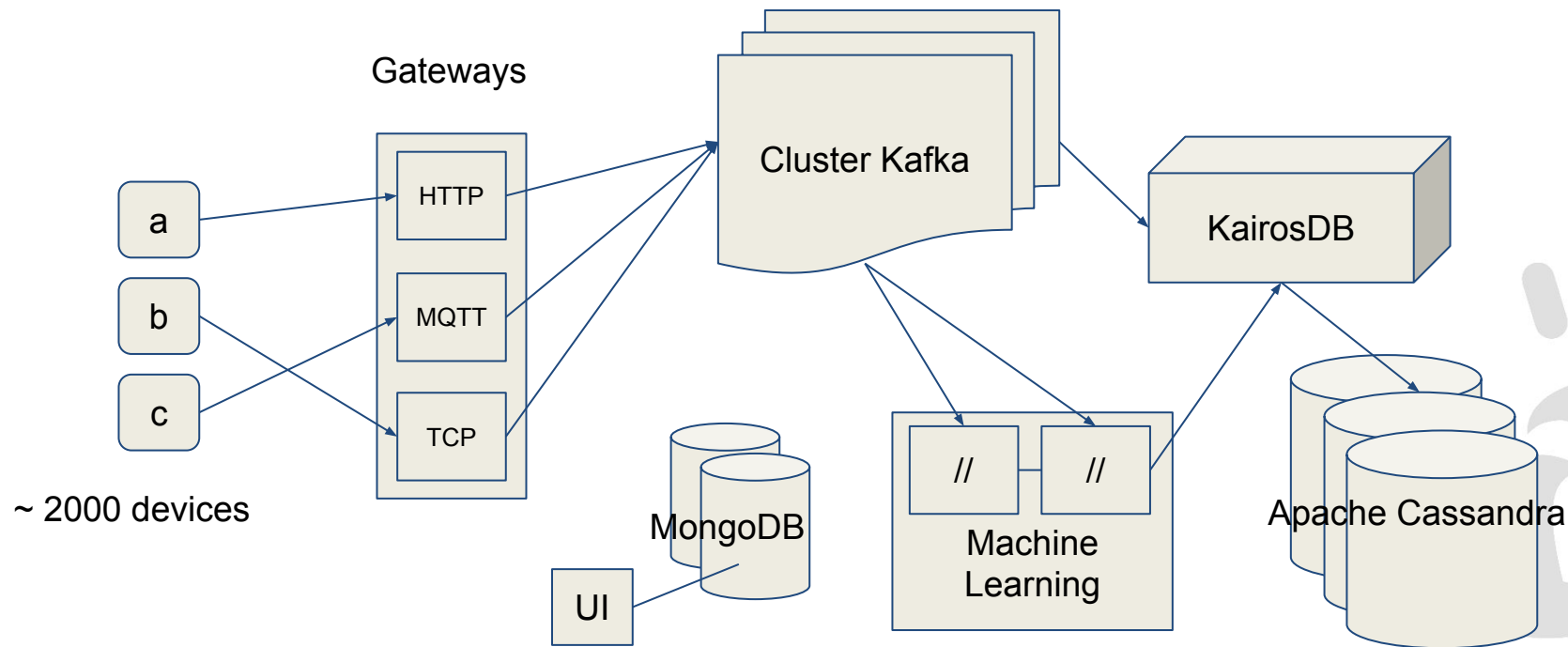


ssantes :

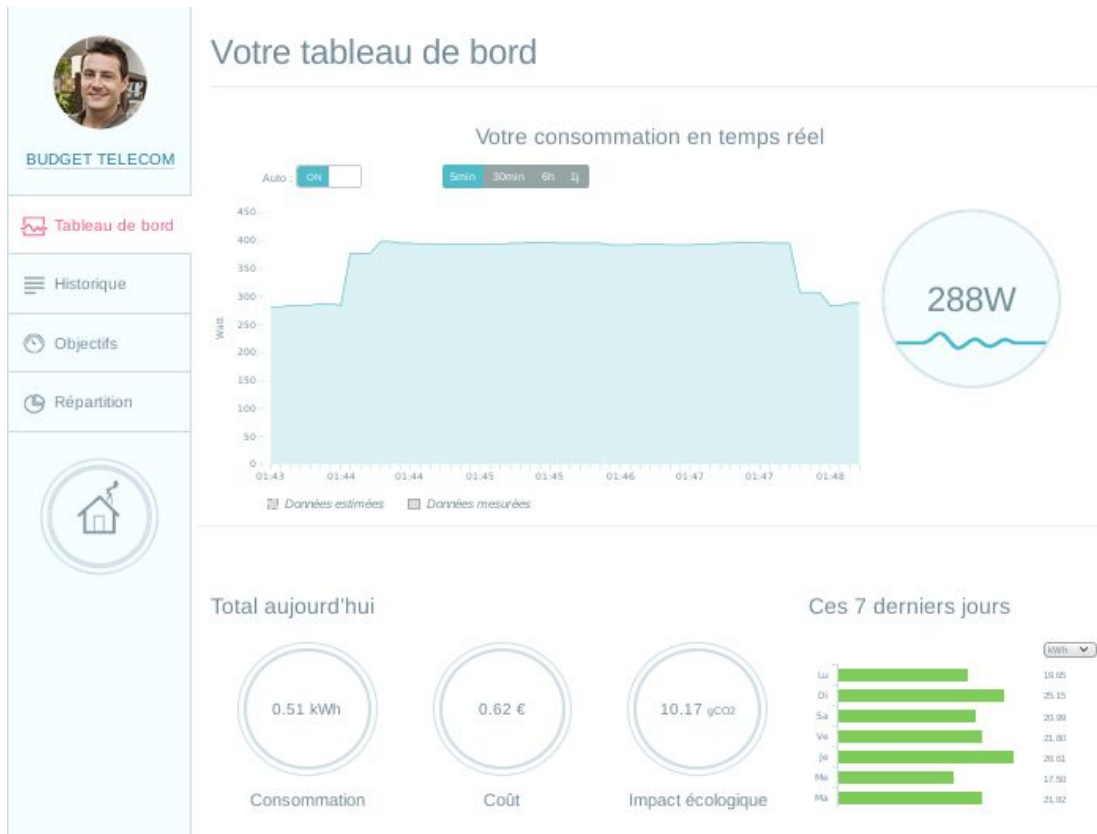
teurs
nnées de

exes et on
vie

2014-2015



On a quand même quelque chose qui tourne..



ça marche et c'est performant !

2015-2016

- On manage Cassandra grâce à DataStax Enterprise
- On passe progressivement la totalité des applications sous Docker
- On migre nos serveurs dédiés chez Google Cloud et on profite de la virtualisation et d'une vraie scalabilité horizontale
- On orchestre la totalité de la plateforme (hors Cassandra) sous Kubernetes

...on essuie d'autres plâtres ;)

- On ma
- On pa
- On mi
virtual
- On ore



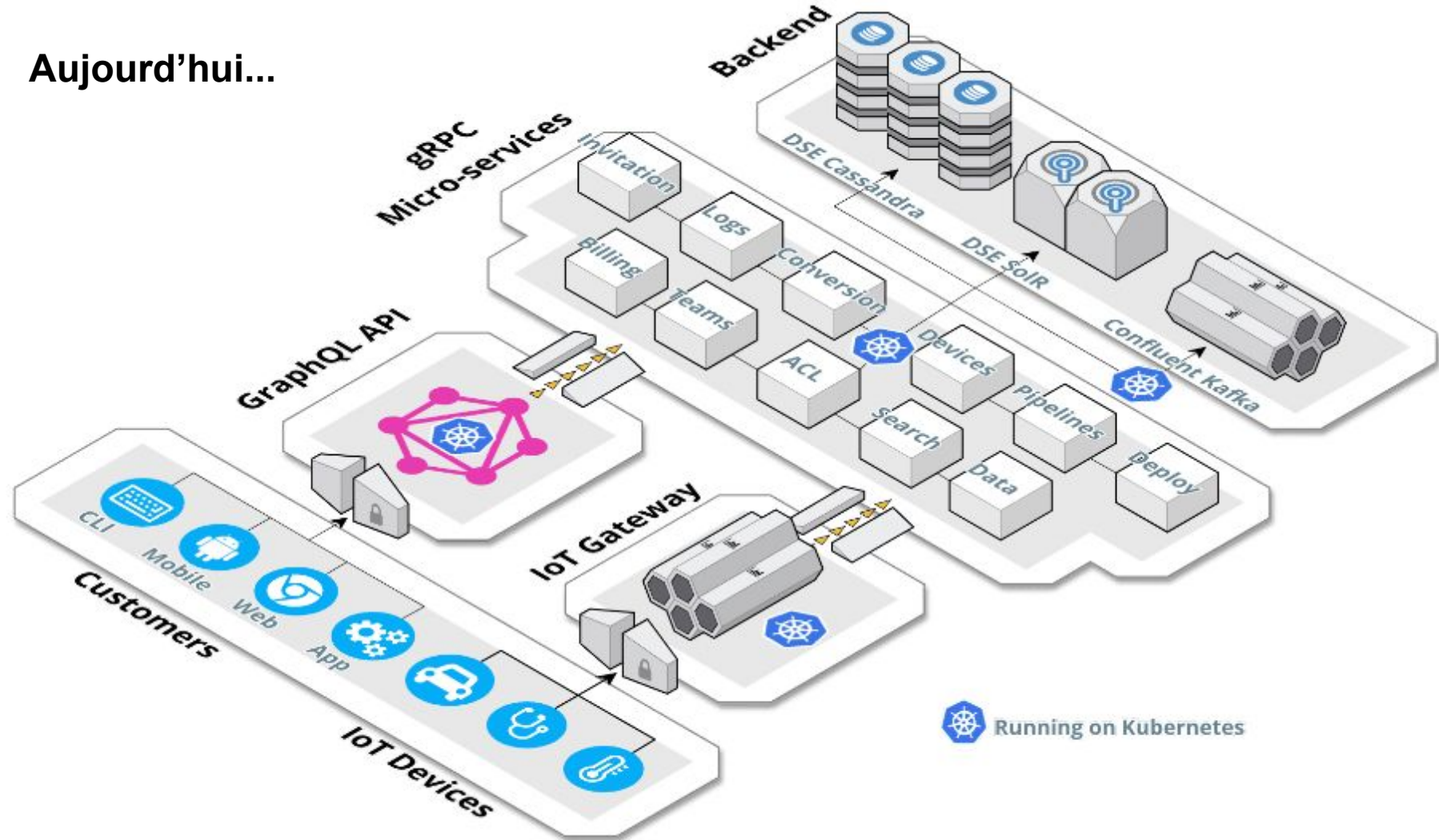
er

e la

ubernetes

tres plâtres ;)

Aujourd'hui...



Pour essayer moins de plâtres..

- Importance de définir un workload cible (batch, temps-réel, r, w, rw)
- Bien choisir le type de base de données NoSQL qui répond au besoin
- Comprendre les implications du modèle de consensus utilisé (PAXOS, Gossip, ...)
- Monitorer monitorer monitorer
- Optimiser très tôt la configuration, mesurer l'impact, REPEAT
- Développer toujours vos applications en prenant en compte que le réseau n'est pas fiable. Préférer les applications Stateless / Crash-only.

Merci de votre attention !
Des questions ?

timelab

