

CubeFS Deep Dive

Возможности единого хранилища для ML,
аналитики и контейнерных платформ

Иван Архипов

SmartData 2026



CubeFS

POSIX

filesystem

HDFS

Hadoop API

S3

object API

Кто говорит?



Иван Архипов

email: me@endevir.ru
tg: @endevir

Чем занимаюсь

Kubernetes

Пишу код в ядро k8s и разрабатываю k8s-операторы для stateful сервисов

Networking

Разрабатываю оверлейную сеть в Яндекс.Облаке

DBMS

В 2023 запустил DBaaS для PostgreSQL в production на десятки тысяч БД в Kubernetes

Storage

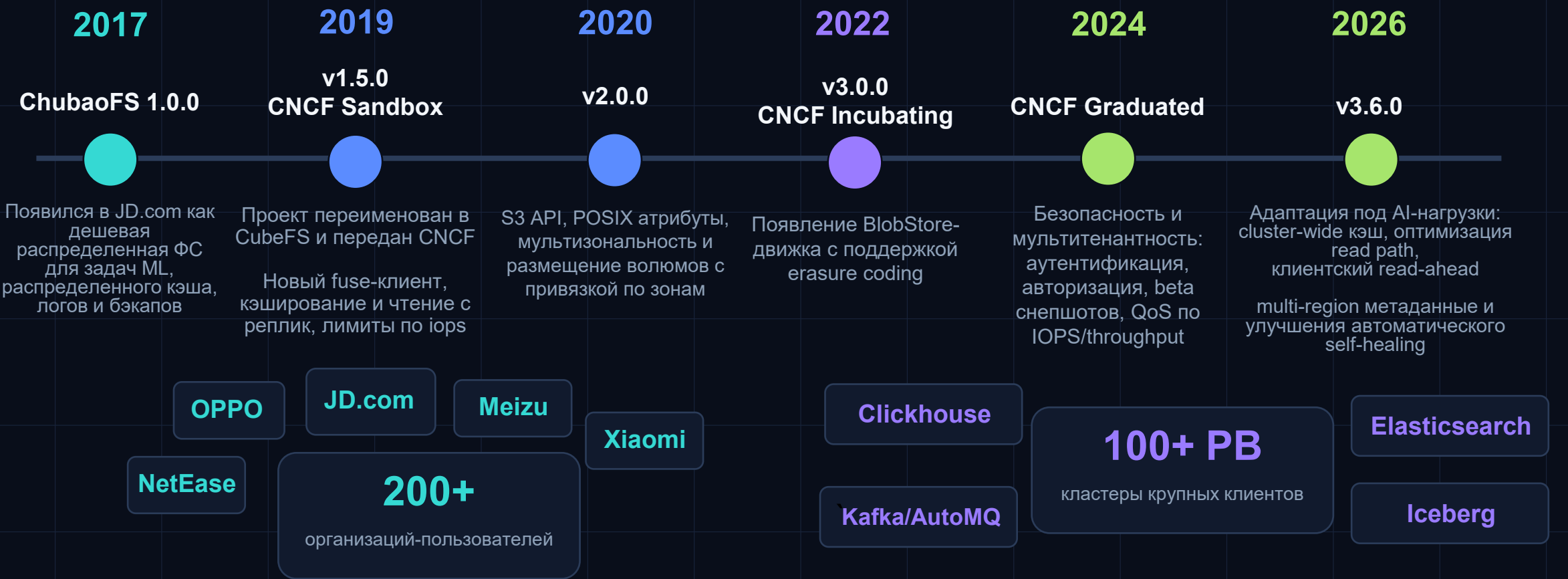
В свободное время собираю и тестирую распределенные хранилки в Kubernetes

Работа с CubeFS – личный проект, поэтому в докладе не будет чисел крупного продакшена, зато будет честный рассказ о том, что болело и как дорабатывали, без NDA :)

Развитие CubeFS



Развитие CubeFS



Устройство CubeFS: общая картина

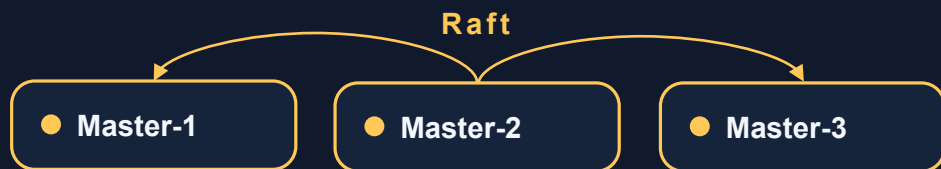
Control Plane

Metadata Plane

Data Plane

Устройство CibeFS: общая картина

Control Plane



Metadata Plane

Data Plane

Устройство CibeFS: общая картина

Control Plane



Metadata Plane

Data Plane

Устройство CibeFS: общая картина

Control Plane



Metadata Plane



Data Plane

Устройство CibeFS: общая картина

Control Plane



Metadata Plane



Data Plane

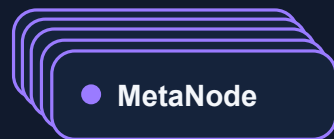


Устройство CibeFS: общая картина

Control Plane



Metadata Plane



Data Plane – Replicated



Устройство CibeFS: общая картина

Control Plane

● Master

● AuthNode

Metadata Plane

● MetaNode

Data Plane – Replicated

● DataNode

Data Plane – Erasure Coding Subsystem

● Access

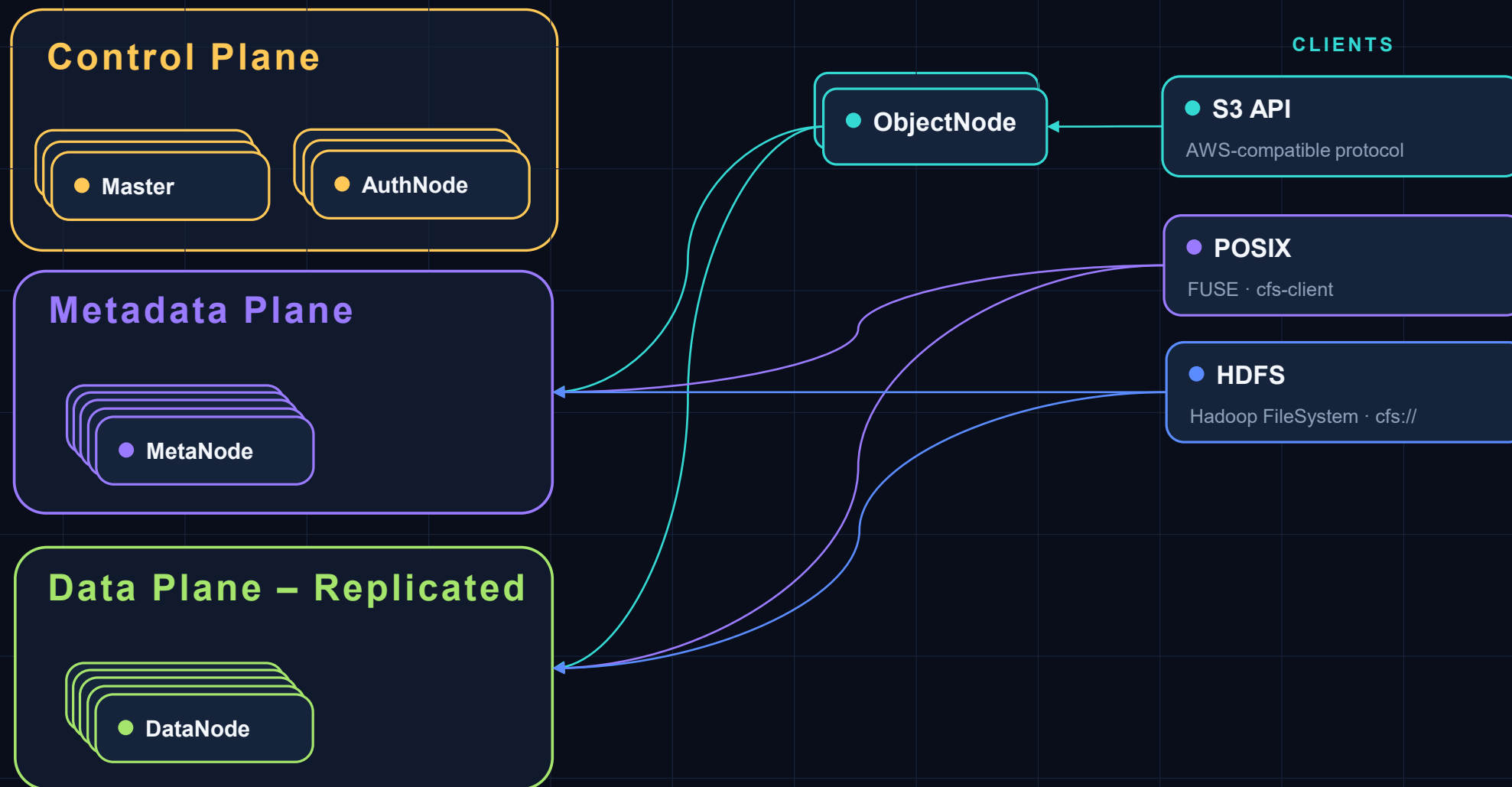
● Proxy

● BlobNode

● ClusterMgr

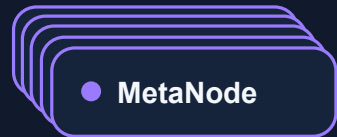
● Scheduler

Устройство CibeFS: общая картина



MetaNode: MetaPartition, Raft и две B-Tree

Metadata Plane



MetaNode: MetaPartition, Raft и две B-Tree

Одна MetaNode содержит много MetaPartition;

Каждая MetaPartition хранит данные только одного Volume.

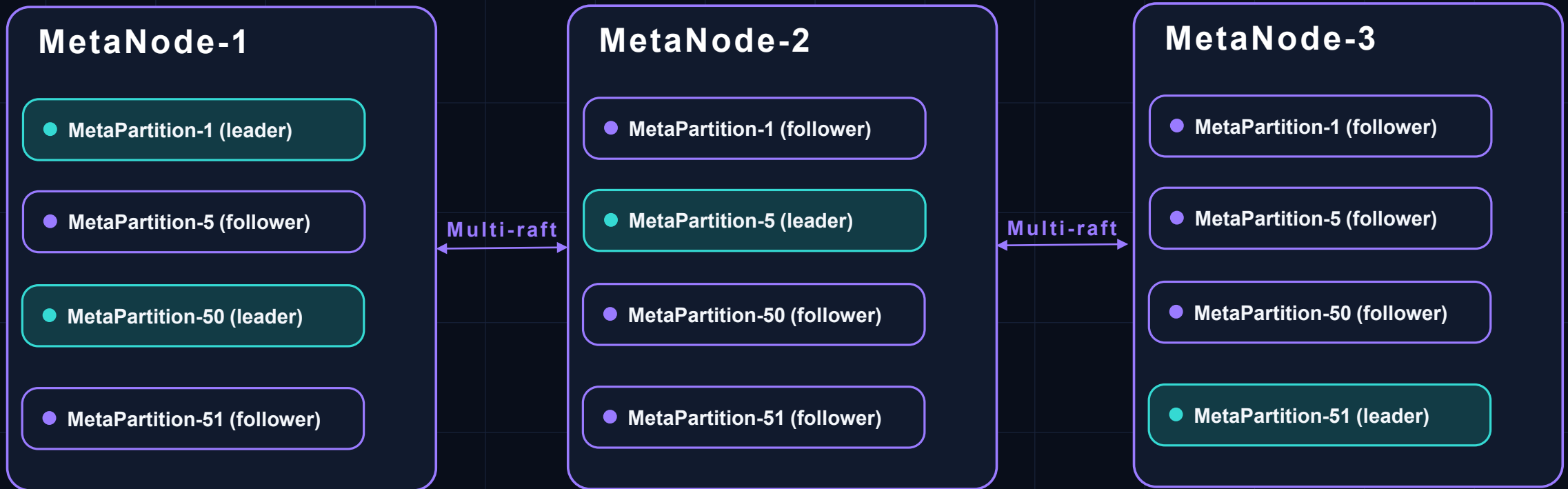
При этом у каждого Volume может быть много MetaPartition, распределенных между MetaNode



MetaNode: MetaPartition, Raft и две B-Tree

MetaPartition реплицируется между разными MetaNode с помощью Raft

У каждой MetaNode могут быть одновременно и партии-лидеры, и партии-фолловеры



MetaNode: MetaPartition, Raft и две B-Tree

Каждый MetaPartition содержит в себе два B-дерева:

- Inode-tree – метаданные об объекте (файле/папке) и размещение его данных
- Dentry-tree – маппинг (parent inode + child name) => child inode, задающий иерархию файловой системы

Inode отвечает на вопросы «что это за объект?» и «где лежат его байты?»

Dentry отвечает за положение объекта в дереве файловой системы.

Dentry · место в дереве

key: (parentId, name)

```
1 // Dentry defines the dentry struct.
2 type Dentry struct {
3     ParentId uint64 // inode of the parent directory
4     Name      string // The name of the current directory entry
5     inode     uint64 // inode of the current file
6     Type      uint32 // File Types and Permissions
7 }
```

Inode – объект ФС

key: inode ID

```
1 type inodeInfo struct {
2     sync.RWMutex
3     inode     uint64 // inode id
4     Type      uint32 // file type and permissions
5     Uid       uint32 // user id
6     Gid       uint32 // group id
7     Size      uint64 // File size
8     Generation uint64 // The version number, incremented for each revision
9     CreateTime int64 // creation time
10    AccessTime int64 // interview time
11    ModifyTime int64 // Change the time
12    LinkTarget []byte // Save the path name of the source file corresponding to the soft link
13    NLink      uint32 // number of hard links
14    Flag       int32 // Whether the mark can be deleted
15    Extents    *SortedExtents // The location information of the file in the three copies
16    ObjExtents *SortedObjExtents // The location information of the file erasure code
17 }
```

MetaNode: MetaPartition, Raft и две B-Tree

Каждый MetaPartition содержит в себе два B-дерева:

- Inode-tree – метаданные об объекте (файле/папке) и размещение его данных
- Dentry-tree – маппинг (parent inode + child name) => child inode, задающий иерархию файловой системы

Inode отвечает на вопросы «что это за объект?» и «где лежат его байты?»

Dentry отвечает за положение объекта в дереве файловой системы.

Чтобы разрешить путь `warehouse/events/day=28/part.parquet`, клиент последовательно ищет dentry каждого элемента дерева, а затем читает inode конечного файла.

Dentry · место в дереве

key: (parentId, name)

```
1 // Dentry defines the dentry struct.
2 type Dentry struct {
3     ParentId uint64 // inode of the parent directory
4     Name      string // The name of the current directory entry
5     inode     uint64 // inode of the current file
6     Type      uint32 // File Types and Permissions
7 }
```

Inode – объект ФС

key: inode ID

```
1 type inodeInfo struct {
2     sync.RWMutex
3     inode     uint64 // inode id
4     Type      uint32 // file type and permissions
5     Uid       uint32 // user id
6     Gid       uint32 // group id
7     Size      uint64 // File size
8     Generation uint64 // The version number, incremented for each revision
9     CreateTime int64 // creation time
10    AccessTime int64 // interview time
11    ModifyTime int64 // Change the time
12    LinkTarget []byte // Save the path name of the source file corresponding to the soft link
13    NLink      uint32 // number of hard links
14    Flag       int32 // Whether the mark can be deleted
15    Extents    *SortedExtents // The location information of the file in the three copies
16    ObjExtents *SortedObjExtents // The location information of the file erasure code
17 }
```

MetaNode: in-memory и сайзинг

Inode- и Dentry-деревья целиком находятся в оперативной памяти, а Raft log и снапшоты позволяют восстановить их после перезапуска.

Ёмкость всего кластера ограничена объёмом оперативной памяти, которую мы резервируем под слой метаданных.

In-memory B-trees

Dentry · место в дереве

key: (parentId, name)

```
1 // Dentry defines the dentry struct.
2 type Dentry struct {
3     ParentId uint64 // inode of the parent directory
4     Name      string // The name of the current directory entry
5     inode     uint64 // inode of the current file
6     Type      uint32 // File Types and Permissions
7 }
```

Inode – объект ФС

key: inode ID

```
1 type inodeInfo struct {
2     sync.RWMutex
3     inode     uint64 // inode id
4     Type      uint32 // file type and permissions
5     Uid       uint32 // user id
6     Gid       uint32 // group id
7     Size      uint64 // File size
8     Generation uint64 // The version number, incremented for each revision
9     CreateTime int64 // creation time
10    AccessTime int64 // interview time
11    ModifyTime int64 // Change the time
12    LinkTarget []byte // Save the path name of the source file corresponding to the soft link
13    NLink      uint32 // number of hard links
14    Flag       int32 // Whether the mark can be deleted
15    Extents    *SortedExtents // The location information of the file in the three copies
16    ObjExtents *SortedObjExtents // The location information of the file erasure code
17 }
```

MetaNode: in-memory и сайзинг

Inode- и Dentry-деревья целиком находятся в оперативной памяти, а Raft log и снапшоты позволяют восстановить их после перезапуска.

Ёмкость всего кластера ограничена объёмом оперативной памяти, которую мы резервируем под слой метаданных.

Оценка для сайзинга - 1ГБ RAM под metanode на 1 млн объектов.

In-memory B-trees

Dentry · место в дереве

key: (parentId, name)

```
1 // Dentry defines the dentry struct.
2 type Dentry struct {
3     ParentId uint64 // inode of the parent directory
4     Name      string // The name of the current directory entry
5     inode     uint64 // inode of the current file
6     Type      uint32 // File Types and Permissions
7 }
```

Inode – объект ФС

key: inode ID

```
1 type inodeInfo struct {
2     sync.RWMutex
3     inode     uint64 // inode id
4     Type      uint32 // file type and permissions
5     Uid       uint32 // user id
6     Gid       uint32 // group id
7     Size      uint64 // File size
8     Generation uint64 // The version number, incremented for each revision
9     CreateTime int64 // creation time
10    AccessTime int64 // interview time
11    ModifyTime int64 // Change the time
12    LinkTarget []byte // Save the path name of the source file corresponding to the soft link
13    NLink      uint32 // number of hard links
14    Flag       int32 // Whether the mark can be deleted
15    Extents    *SortedExtents // The location information of the file in the three copies
16    ObjExtents *SortedObjExtents // The location information of the file erasure code
17 }
```

MetaNode: шардирование MetaPartition

MetaPartition делят метаданные по диапазонам inode ID автоматически по достижении порога по RAM



MetaNode: шардирование MetaPartition

MetaPartition делят метаданные по диапазонам inode ID автоматически по достижении порога по RAM



Что шардируется

Dentry records

Запись индексируется парой parentId + name и хранится с parent.

Inode records

Новые inode попадают в следующий диапазон MetaPartition.

Что не шардируется

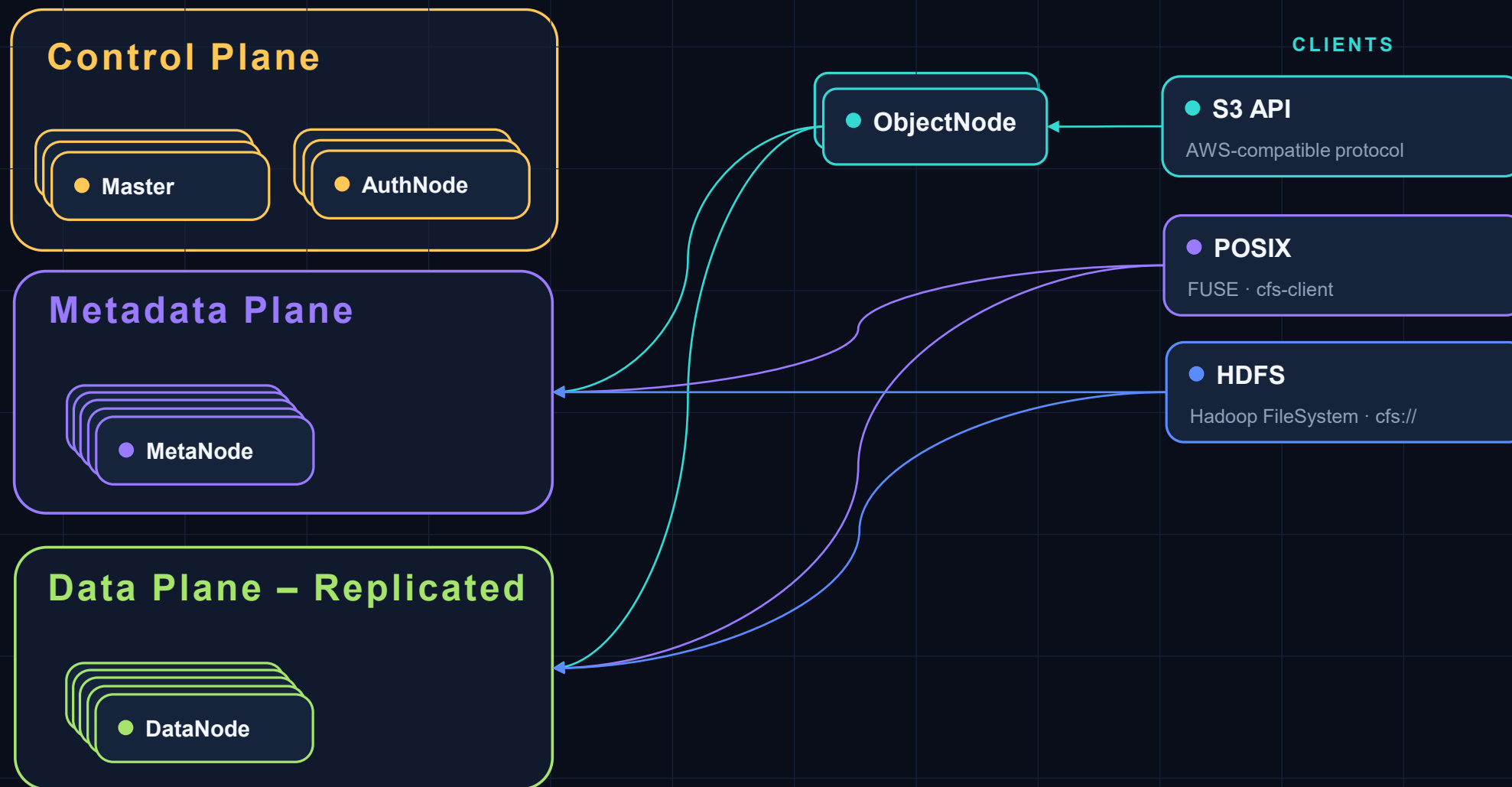
Одна Inode - неделима

Следствие: файлы сверхбольших размеров будут дополнительно нагружать MetaNode за счёт раздутой Inode

Одна Dentry - неделима

Следствие: нельзя создавать папки на >20млн. файлов

Устройство CibeFS: общая картина



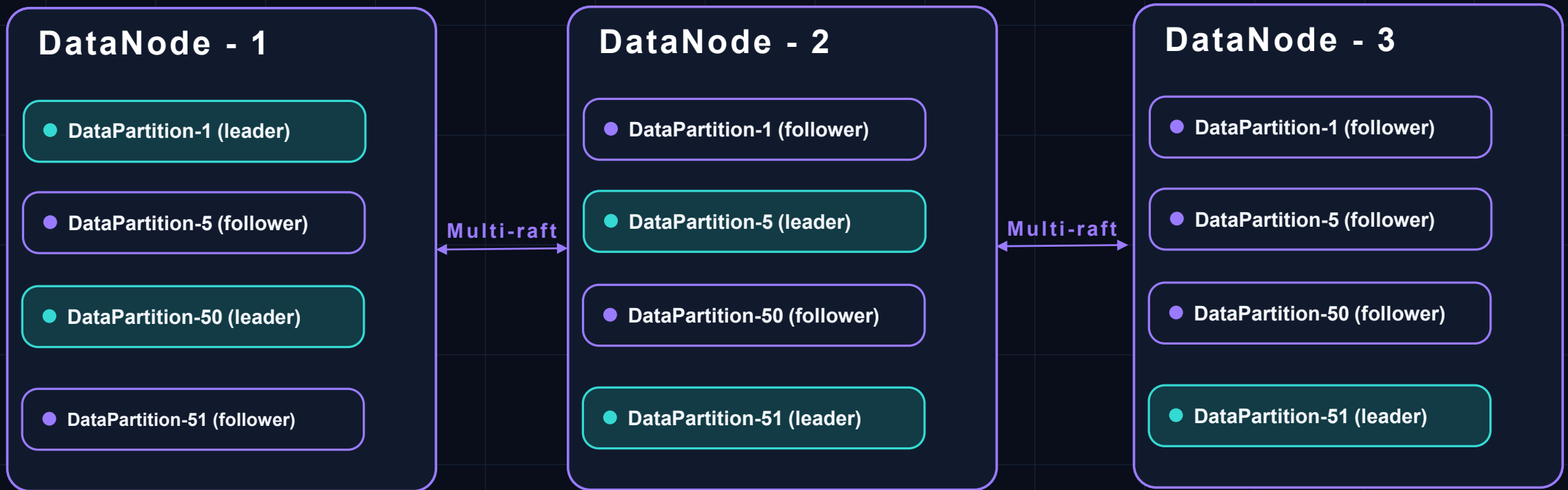
Устройство CubeFS: общая картина

Data Plane – Replicated



Data plane: Replica subsystem

DataPartition реплицируется между разными DataNode с помощью Raft
У каждой DataNode могут быть одновременно и партии-лидеры, и партии-фолловеры

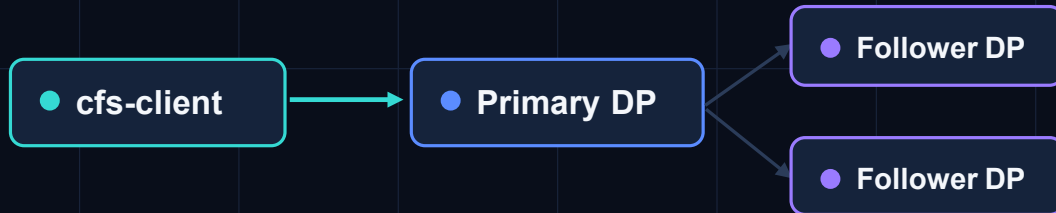


~120 GB

типовой DataPartition

Data plane: Replica subsystem

Клиент сам выбирает DataPartition для записи, DataPartition рассылает обновления на фолловеров

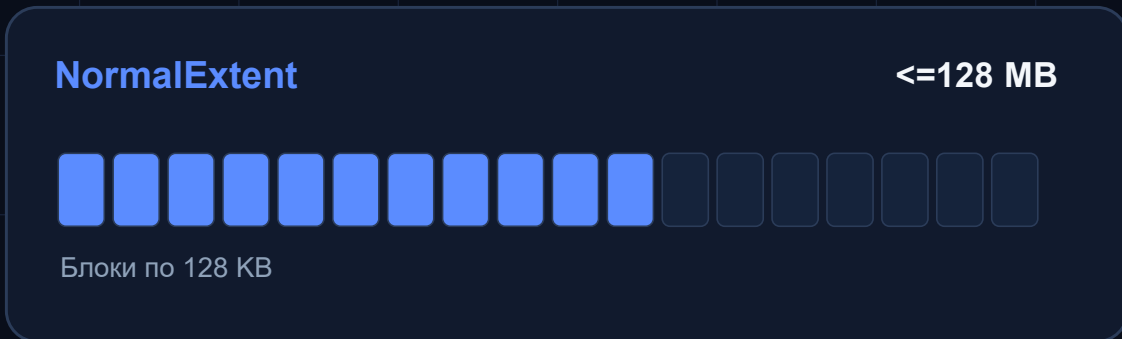
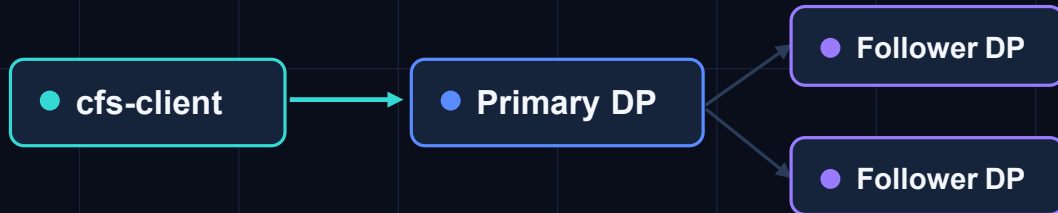


При последовательной записи новых данных Primary DP синхронно рассылает изменения во все follower DP. Если это не удалось – клиент повторяет попытку уже в другую партицию.

При перезаписи существующих данных используется ослабленное требование – достаточно кворума из DP, т.к. возможности ретраиться в другие DP у клиента нет.

Data plane: Replica subsystem

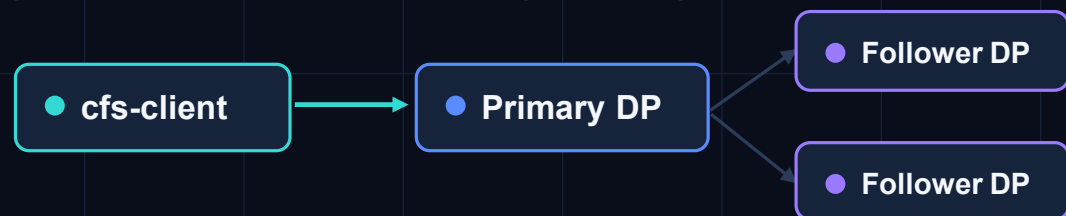
Клиент сам выбирает DataPartition для записи, DataPartition рассылает обновления на фолловеров



Обычные и крупные объекты записываются в файлы NormalExtent блоками по 128КБ. Файл NormalExtent ограничен 1024 блоками.

Data plane: Replica subsystem

Клиент сам выбирает DataPartition для записи, DataPartition рассылает обновления на фолловеров



NormalExtent

<=128 MB



Блоки по 128 KB

Обычные и крупные объекты записываются в файлы NormalExtent блоками по 128KB. Файл NormalExtent ограничен 1024 блоками.

TinyExtent

64 файла на partition, размер файла – до 4 TiB



Блоки по 4 KB

Для небольших объектов есть специализированный путь — файлы TinyExtent.

CubeFS агрегирует много маленьких объектов в общие append-файлы.

Data plane: Replica subsystem – особенности

Простота и надёжность

Нижележащее хранилище – привычная файловая система (ext4, xfs, btrfs, etc.)

Поддержка in-place перезаписи объектов

Вместо copy-on-write на весь блок или файл

Прозрачное multi-tier хранилище

Данные автоматически перемещаются в более дешёвые хранилища в соответствии с политиками жизненного цикла

Управляемое горизонтальное масштабирование чтения-записи

Для каждого Volume можем задавать, сколько активных Meta- и Data-партиций держать

Data plane: Replica subsystem – особенности

Простота и надёжность

Нижележащее хранилище – привычная файловая система (ext4, xfs, btrfs, etc.)

Поддержка in-place перезаписи объектов

Вместо copy-on-write на весь блок или файл

Прозрачное multi-tier хранилище

Данные автоматически перемещаются в более дешевые хранилища в соответствии с политиками жизненного цикла

Управляемое горизонтальное масштабирование чтения-записи

Для каждого Volume можем задавать, сколько активных Meta- и Data-партиций держать

Нет полноценной поддержки снапшотов

В subefs реализован механизм снапшотов, но на данный момент апстрим не поддерживает multi-tier volume'ы и функциональность отключена по-умолчанию

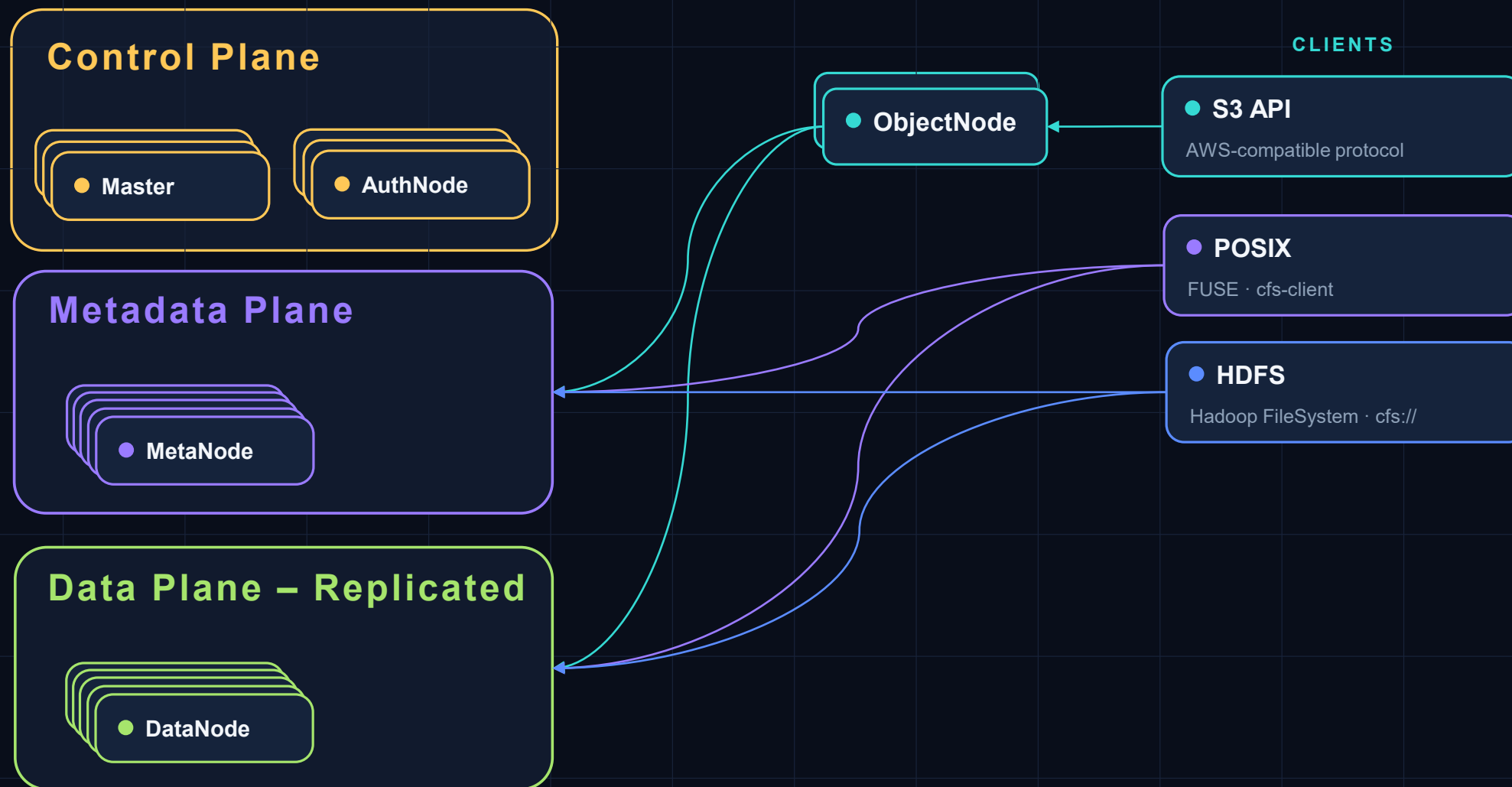
Асинхронная проверка размера партиций

Заполненные партиции не сразу переходят в RO по достижении порога в 120 ГБ. Рекомендуется добавлять 5-10% запаса при расчетах сайзинга

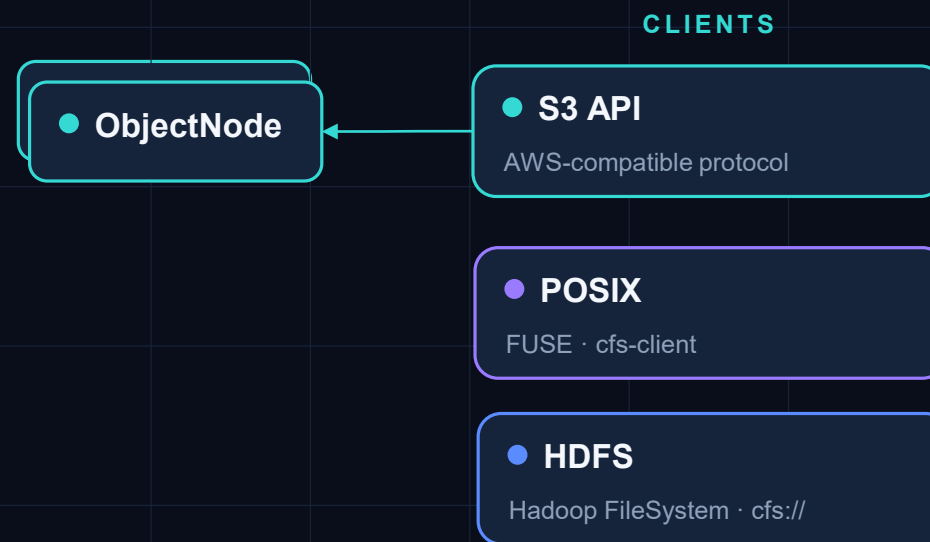
Отсутствует авто-детекция bitrot

Для блоков в Data-партициях сохраняются CRC-чексуммы, но нет аналогов Deep Scrub ceph'a. Автоматика проверяет на I/O ошибки диска и расхождения в размерах extent'ов (неполную запись)

Устройство CibeFS: общая картина



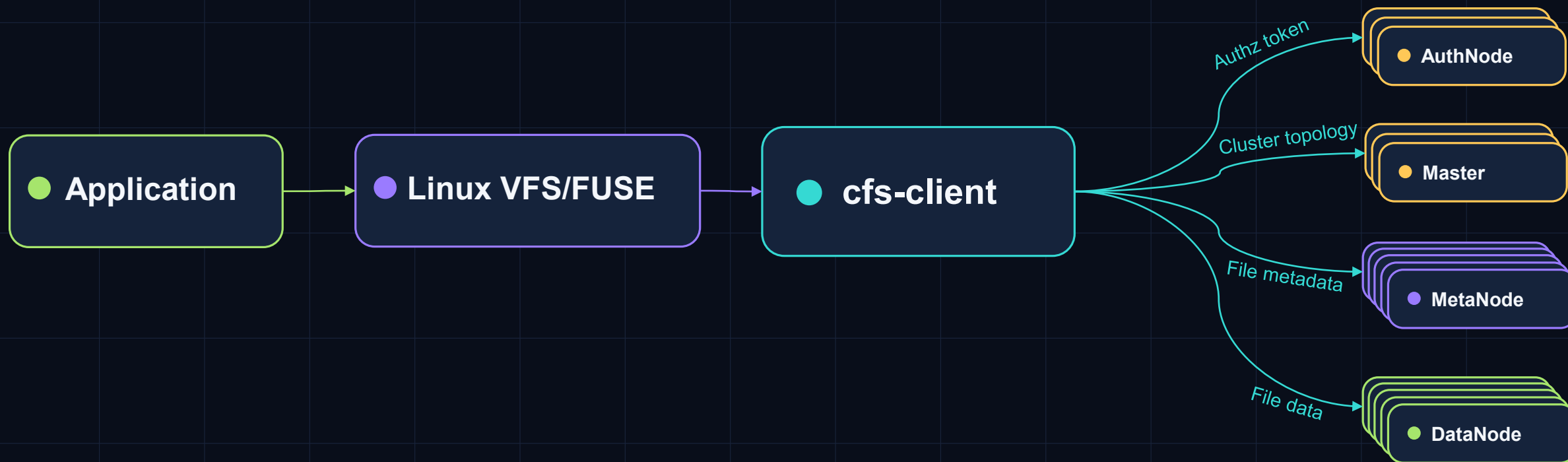
Устройство CibeFS: общая картина



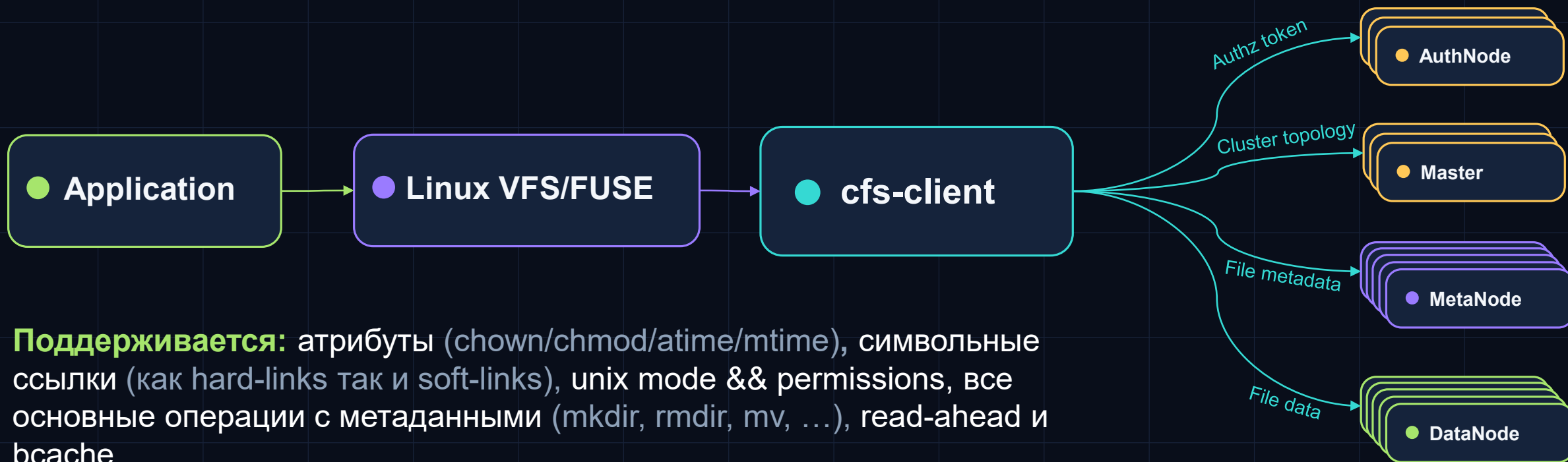
POSIX: возможности и ограничения файловой системы



POSIX: возможности и ограничения файловой системы



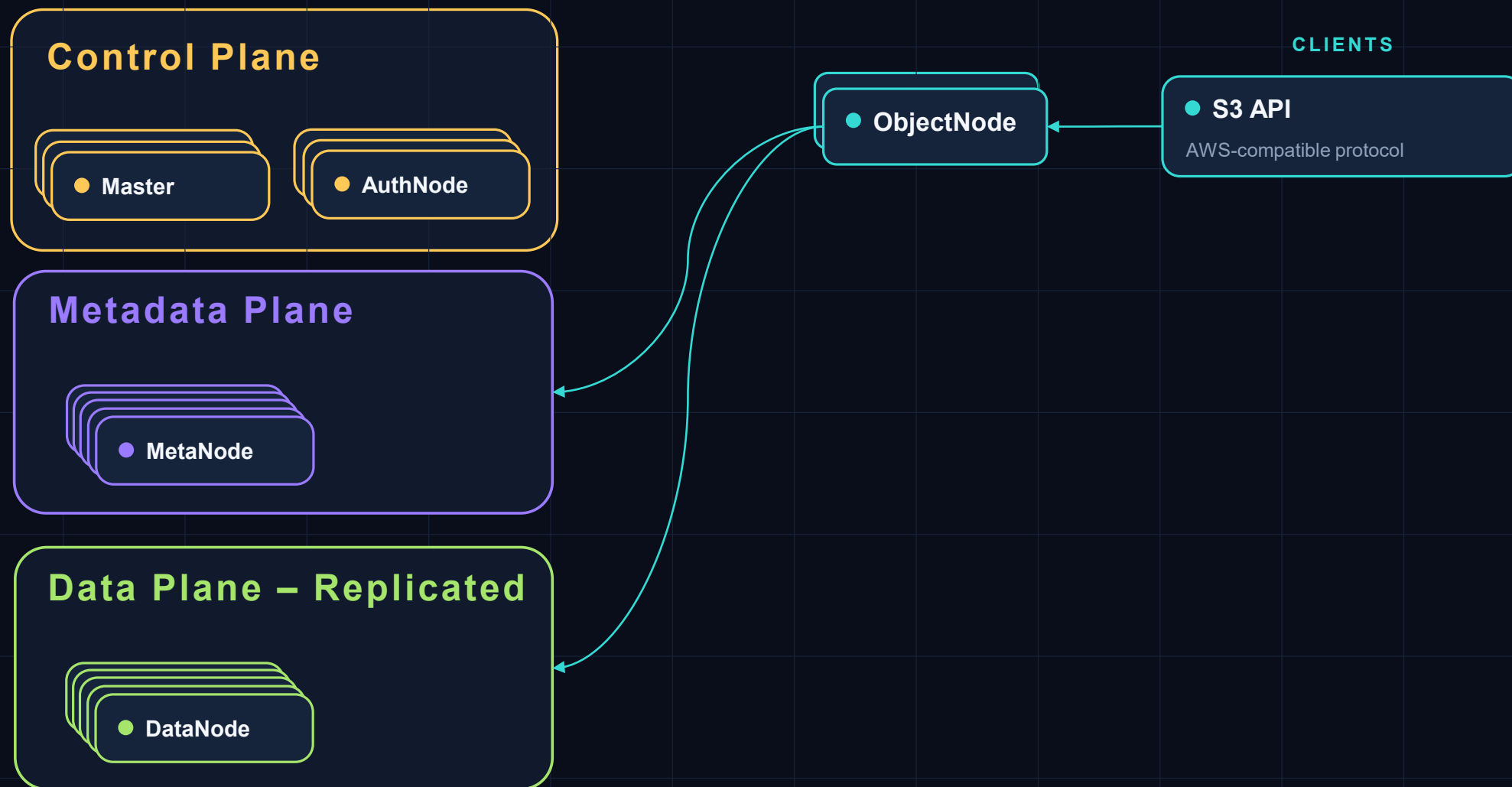
POSIX: возможности и ограничения файловой системы



Поддерживается: атрибуты (chown/chmod/atime/mtime), символные ссылки (как hard-links так и soft-links), unix mode & permissions, все основные операции с метаданными (mkdir, rmdir, mv, ...), read-ahead и bcache

Не поддерживается: flock и распределенные блокировки на запись, fsnotify

S3 API: возможности и ограничения ObjectNode



S3 API: возможности и ограничения ObjectNode

PUT `events/2026/08/part.parquet`

events

2026

08

part.parquet

Ключи S3 объектов становятся реальными папками и файлами во внутреннем представлении

При создании объекта рекурсивно создаются все недостающие папки в дереве

При удалении объекта в штатном CubeFS удаляется только непосредственно файл. Папки остаются на месте.

У себя реализовали опциональный механизм рекурсивной очистки при удалениях объектов через S3 API, получилось дешево.

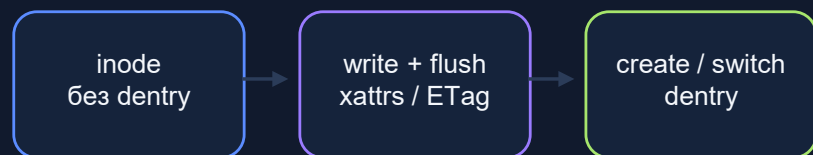
S3 API: возможности и ограничения ObjectNode

PUT `events/2026/08/part.parquet`



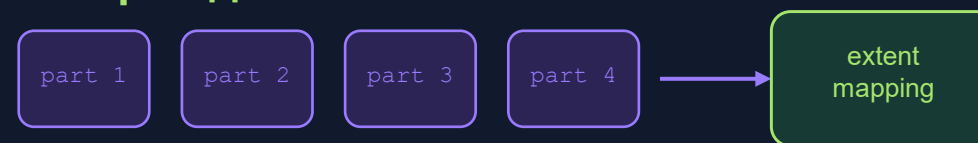
Ключи S3 объектов становятся реальными папками и файлами во внутреннем представлении

PUT большого объекта – атомарный



Читатель не увидит частично загруженный объект

Multipart – атомарный и параллельный поверх одной inode



При завершении ObjectNode предоставляет правильные офсеты загруженных данных и атомарно публикует inode аналогично с единичным PUT

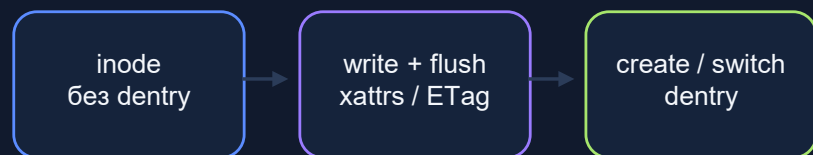
S3 API: возможности и ограничения ObjectNode

PUT events/2026/08/part.parquet



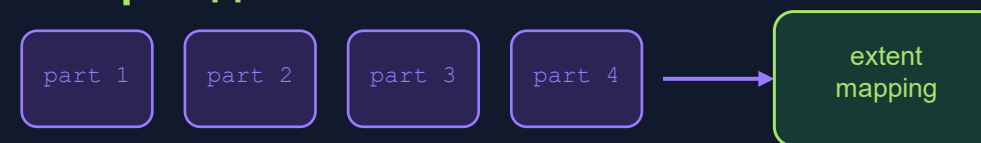
Ключи S3 объектов становятся реальными папками и файлами во внутреннем представлении

PUT большого объекта – атомарный



Читатель не увидит частично загруженный объект

Multipart – атомарный и параллельный поверх одной inode



При завершении ObjectNode предоставляет правильные офсеты загруженных данных и атомарно публикует inode аналогично с единичным PUT

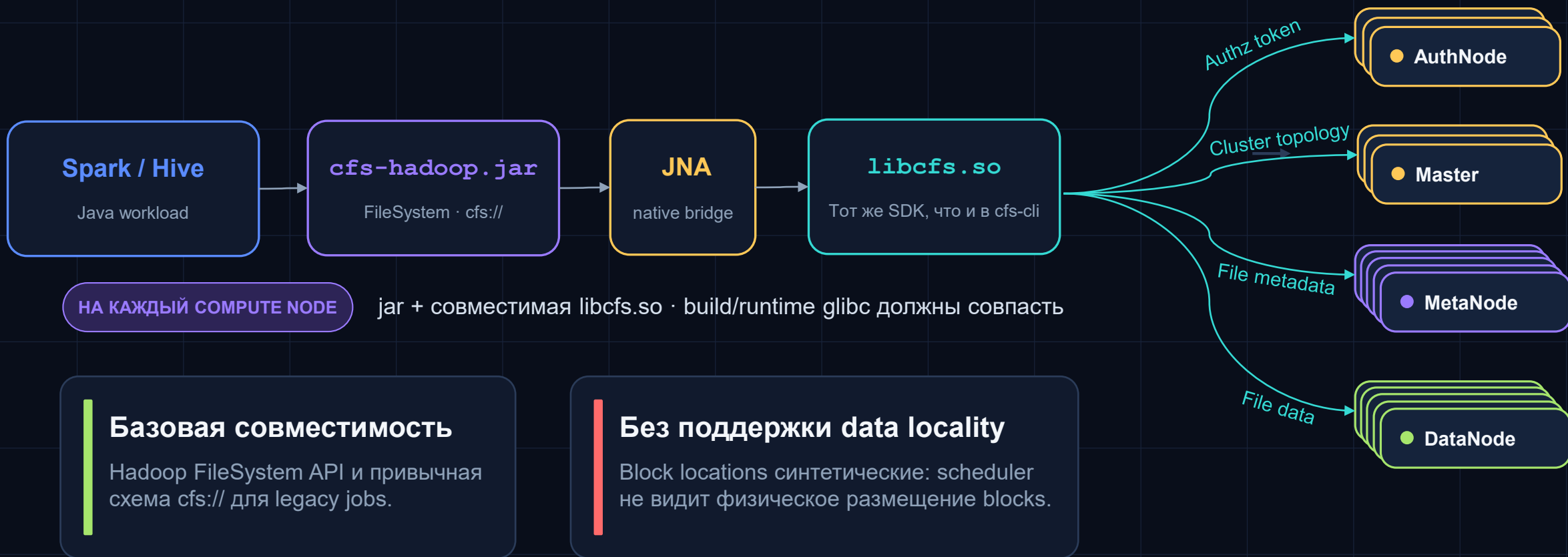
Ограничения, которые слишком сложно преодолеть:

- Версионирование объектов
- Невозможность создать объект a/b и затем создать объект a/b/c
- Тяжёлый ListObjects, который требует запроса по множеству MetaNode

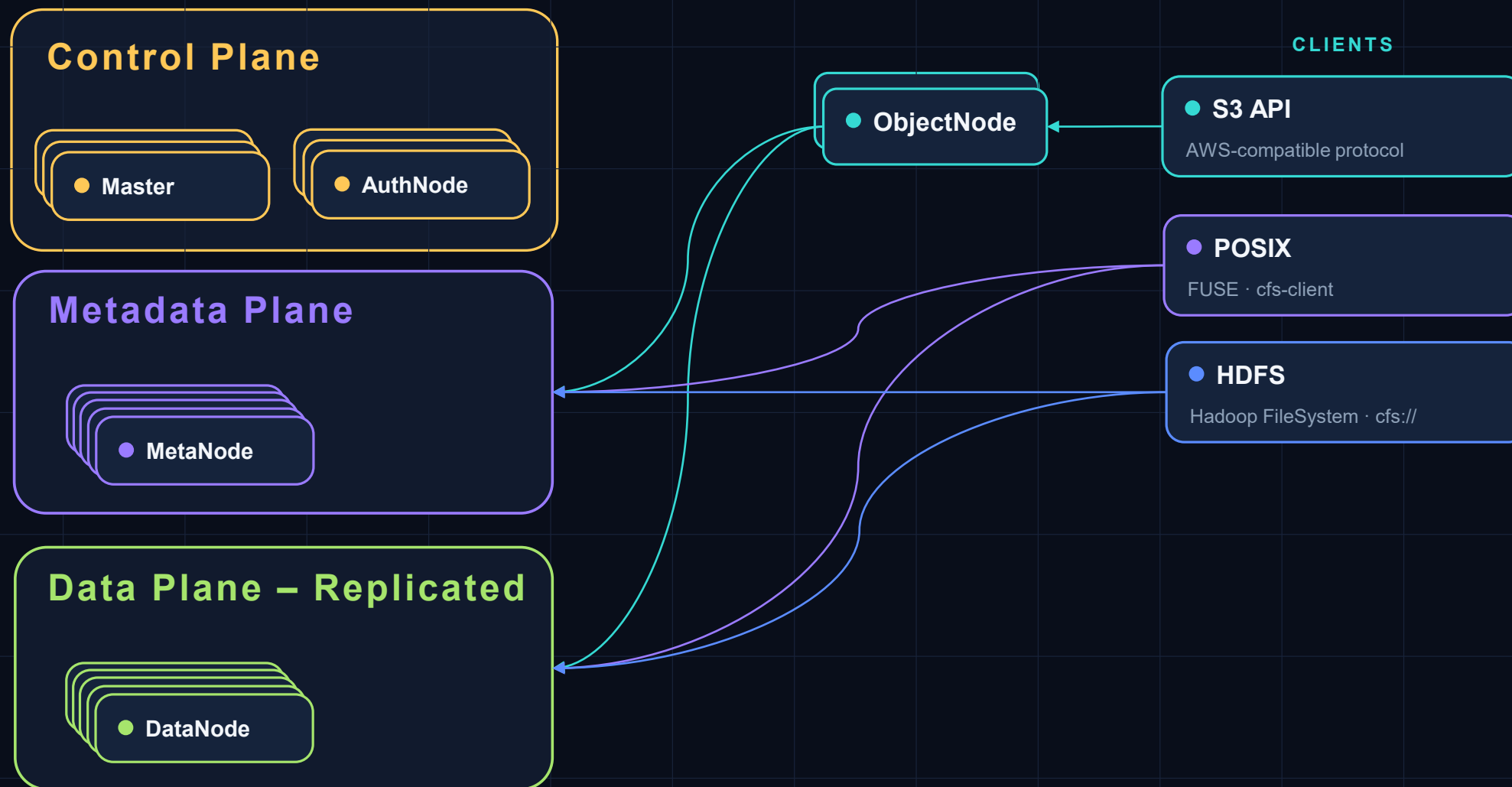
HDFS adapter: возможности и ограничения



HDFS adapter: возможности и ограничения



Устройство CibeFS: общая картина



Опыт и нюансы эксплуатации

- Нет Backfill на добавлении новых нод, кластер может оставаться неравномерно заполненным
- UI де-факто заброшен, необходимо писать свой
- Документация и логи достаточно сложные для восприятия
- Меньше удобства в управлении топологией, чем у Serp
- Развертывание в k8s helm-чартом == невозможность управлять топологией кластера в рантайме)
- Мало lifecycle-политик (не до конца реализован Expiration по времени)

Опыт и нюансы эксплуатации

- Нет Backfill на добавлении новых нод, кластер может оставаться неравномерно заполненным
- UI де-факто заброшен, необходимо писать свой
- Документация и логи достаточно сложные для восприятия
- Меньше удобства в управлении топологией, чем у Serp
- Развертывание в k8s helm-чартом == невозможность управлять топологией кластера в рантайме)
- Мало lifecycle-политик (не до конца реализован Expiration по времени)
- Написан на Go и легко дорабатывается
- Экономичный по ресурсам, легко разворачивается в Kind
- Нет Backfill на добавлении новых нод => нет больших перемещений данных при вводе :)

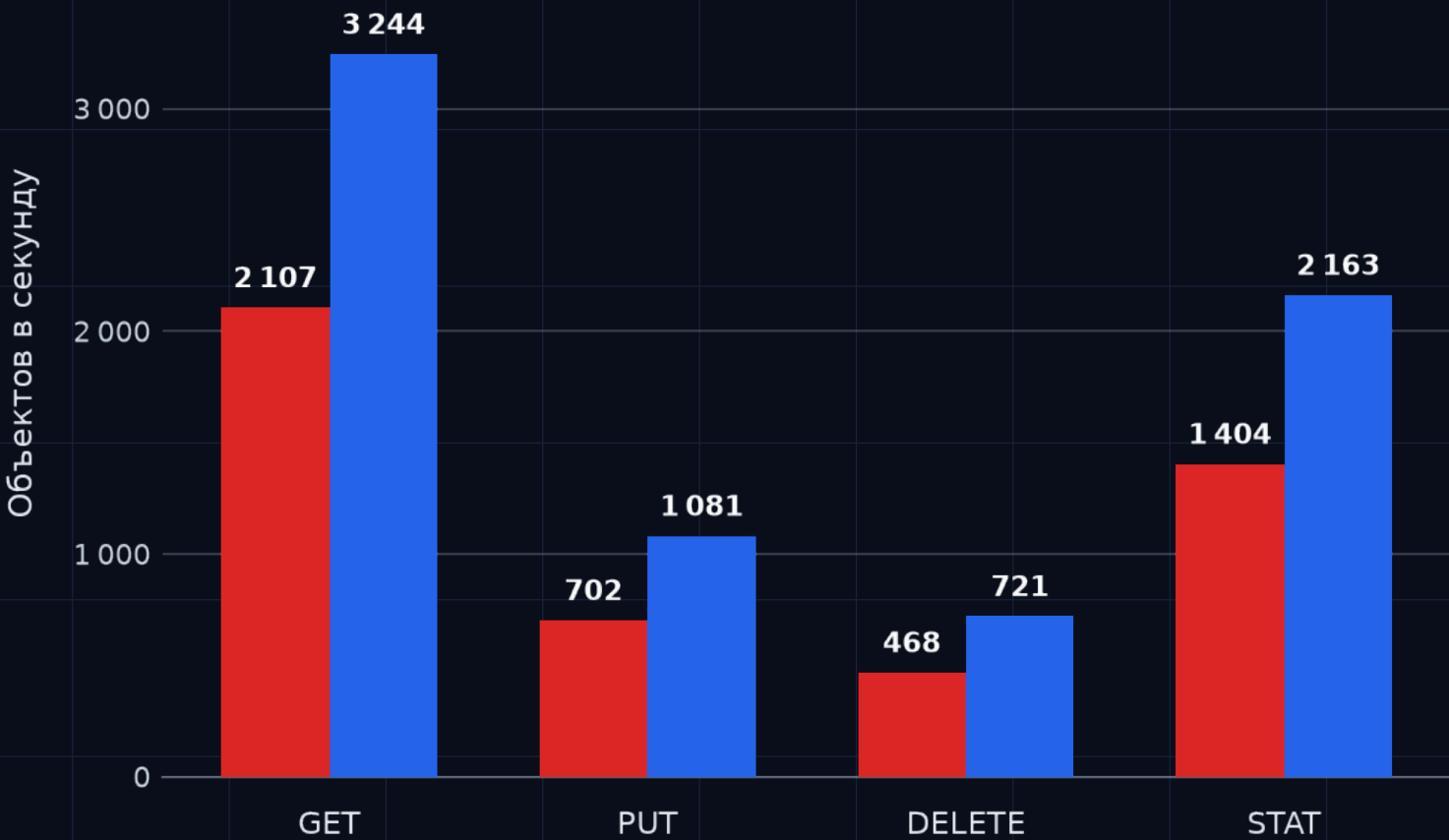
CubeFS и Ceph: сравнительная производительность

Средняя производительность по типам операций · больше — лучше

Ceph CubeFS

Файловая система · fio randrw 70/30, 4 KiB

S3 API · Warp mixed, 64 KiB



Тесты проводились на dev-кластере из 4х машин с 16 vCPU, 32 GB RAM, 1000 QLC SSD, 10 Gb/s ethernet
Перепроверяйте результаты на своих мощностях и конфигурациях!

CubeFS сейчас

- Быстрая работа с метаданными на миллиардах объектов
- Одновременная поддержка S3, HDFS и POSIX API над одними и теми же данными
- Поддержка частичной перезаписи объектов
- Многоуровневое кэширование
- Прозрачное гибридное хранилище

Но нужен ресурс разработки/SRE на доработку для удобной эксплуатации!