



Resource Management Systems for Complex & Non-Predictably Evolving Applications

Cristian Klein & Christian Perez

AVALON Team, INRIA/LIP, ENS Lyon, France

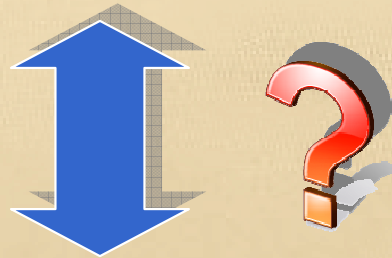
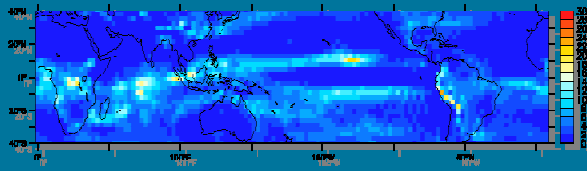
Intl. Research Workshop Advanced HPC Systems

Cetraro, Italy, June 28th, 2011



Avalon Research Topics

Applications



Hybrid

**Cluster
(GPU/MC/...)**

**Grids
(EGEE)**

Heterogeneity

Large
scale

**Super-
computers
(Exascale)**

Clouds

Elasticity

Scientific applications

- Complexity
 - Code coupling
- Code and language heterogeneity
- Performance needs
 - Computation
 - Data

Objectives

- Expressiveness simplicity
- Application portability
- Resource specific optimizations
- Elastic resource management

Content

- *Context*

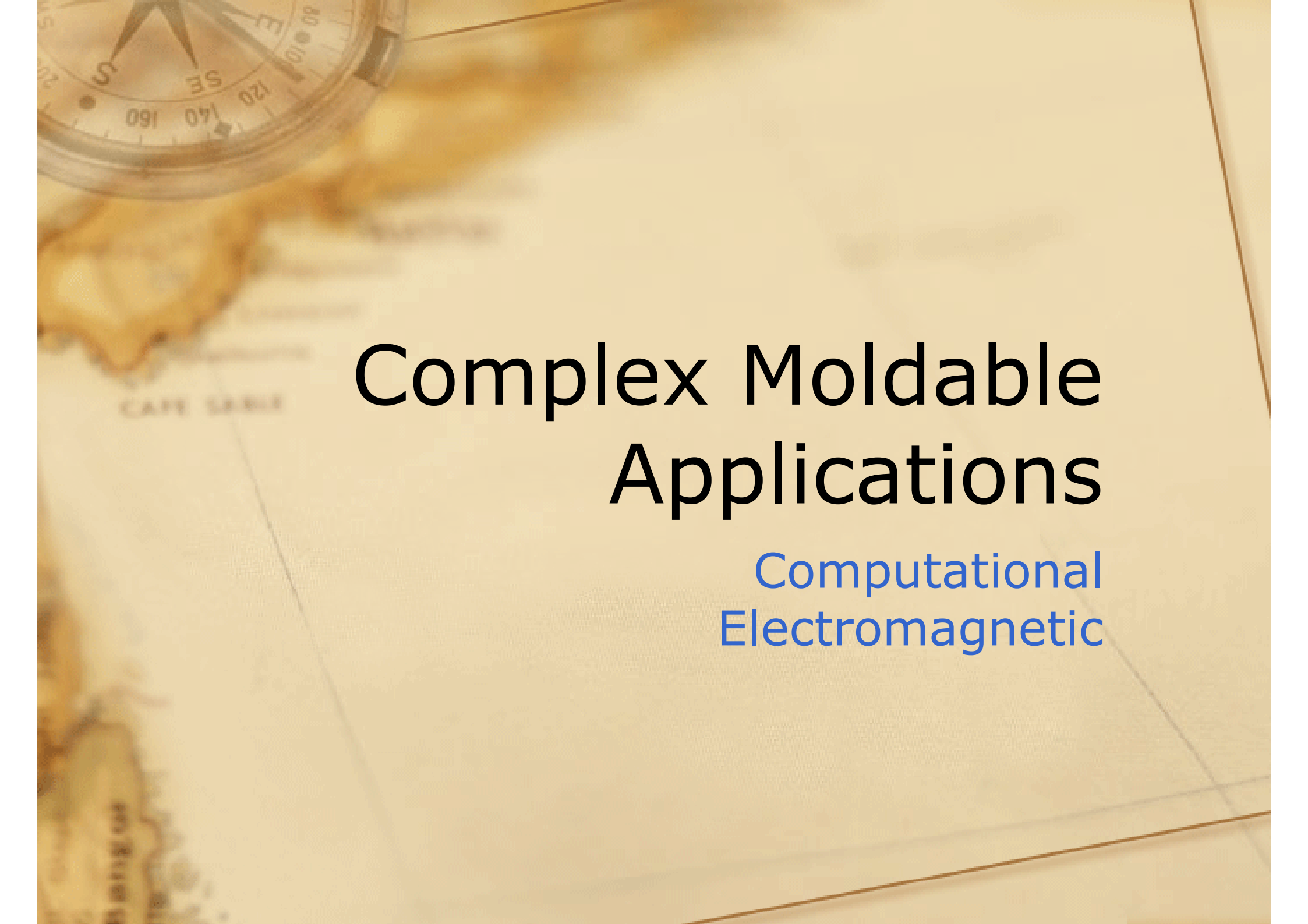
- Complex moldable applications

- Example (Computational Electromagnetic)
- Limitations of existing RMS
- CooRMv1 – Model & Experiments

- Non predictably evolving applications

- Example (AMR)
- Limitations of existing RMS
- CooRMv2 – Model & Experiments

- Conclusion and future works

The background of the slide is a vintage map with a compass rose in the top left corner. The map is aged and yellowed, with various geographical features and labels. The compass rose shows cardinal and intercardinal directions (N, NE, E, SE, S, SW, W, NW) and degree markings. The text "Complex Moldable Applications" is centered on the map in a large, black, sans-serif font.

Complex Moldable Applications

Computational
Electromagnetic

Computational Electromagnetic

- From ANR DISCOGRID

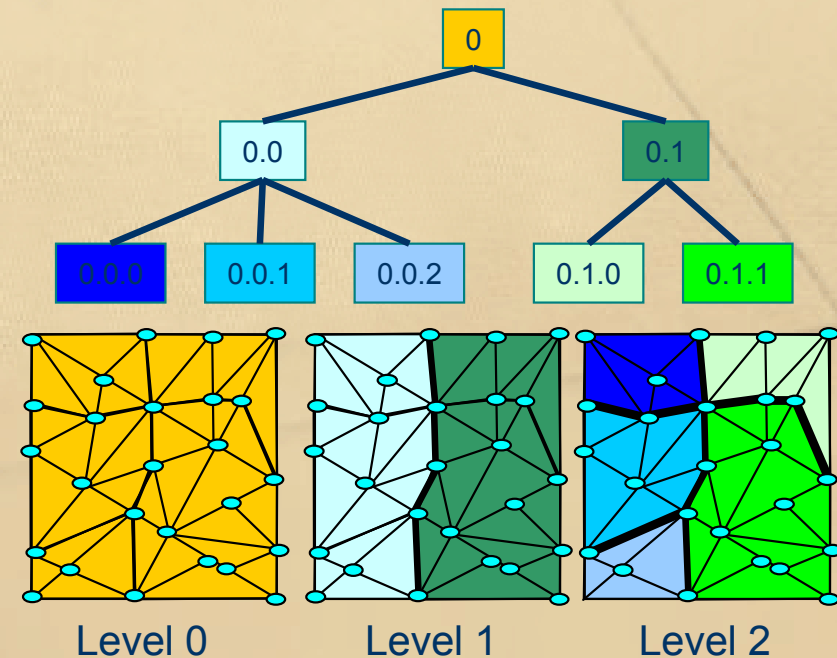
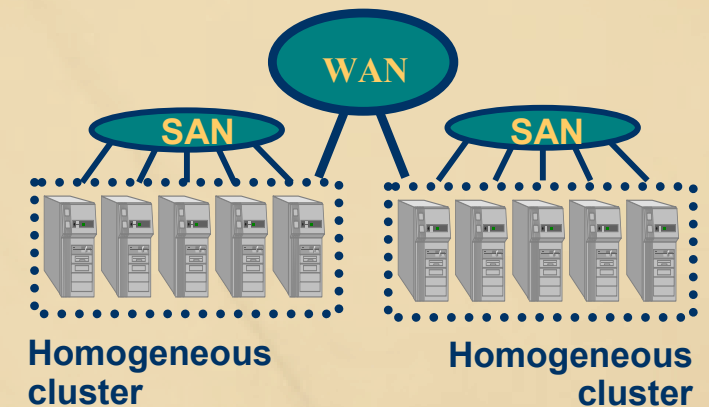
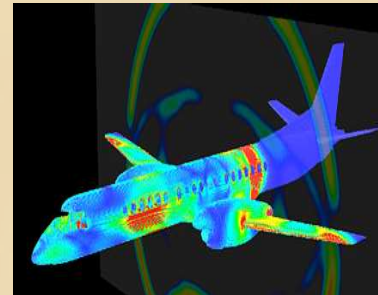
- ANR-05-CIGC-005

- Resource model

- Hierarchical machines
 - Federation of clusters

- Application model

- Set of MPI-based codes
 - How many groups?
 - Size of each group?



Application Performance Model

Inputs

- Number of tetrahedra: n_t
- Clusters: n_H^i , B_H^i , B_C^i , α_j^i
- Maximum inter-cluster (WAN) latency: l_{WAN}

One iteration

$$n_{tl}^{(i)} = n_t \cdot \frac{\sum_j 1/\alpha_j^{(i)}}{\sum_{k \in C} \left(n_H^{(k)} \cdot \sum_j 1/\alpha_j^{(k)} \right)}$$

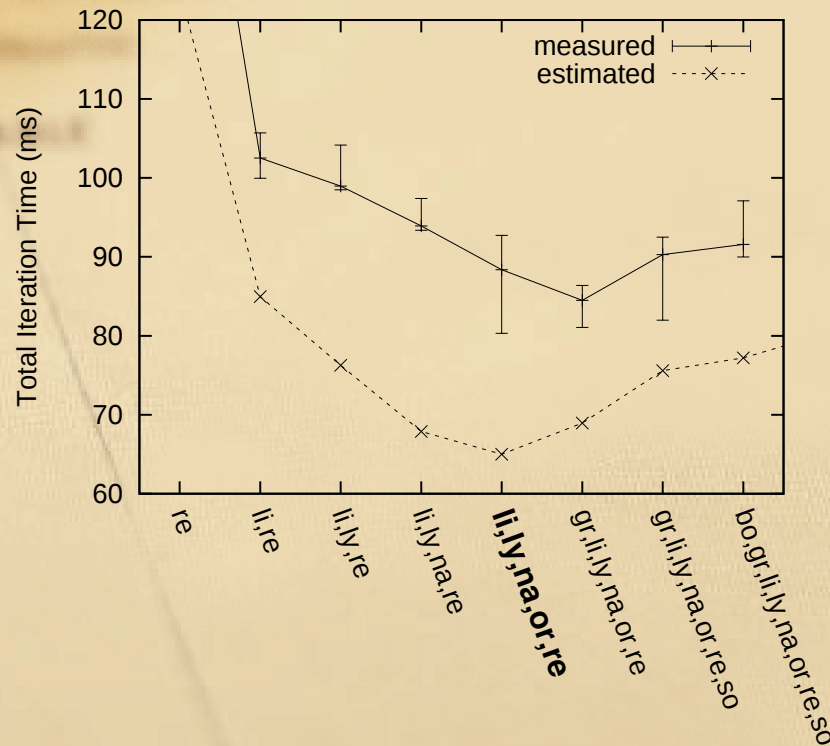
$$t = \sum_{j \in \text{non-overlap}} \max_{i \in C} \left(\alpha_j^{(i)} \cdot n_{tl}^{(i)} \right)$$

$$+ \sum_{j \in \text{overlap}} \max_{i \in C} \left(\alpha_j^{(i)} \cdot n_{tl}^{(i)}, l_{WAN} + \frac{s_U}{\min(B_C^{(j)}, n_H^{(j)} \cdot B_H^{(j)})} \cdot \beta_C \cdot \left(\frac{n_t}{n_C} \right)^{2/3} + \frac{s_U}{B_H^{(i)}} \cdot \beta_H \cdot \left(\frac{n_t}{\sum_{i \in C} n_H^{(i)}} \right)^{2/3} \right)$$

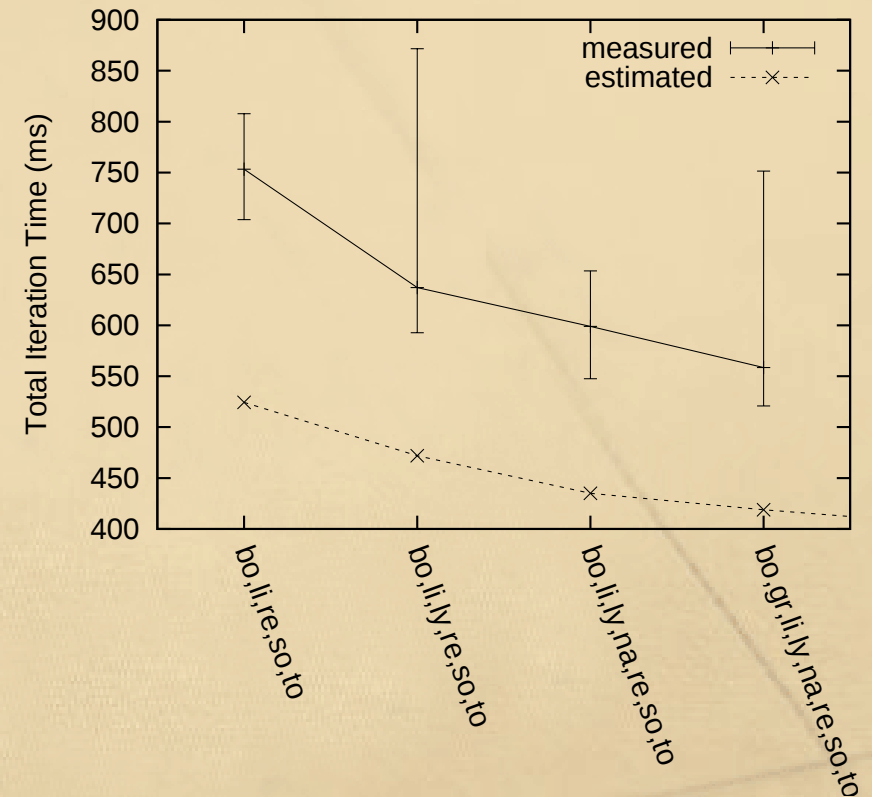
$$+ n_{AR} \cdot l_{WAN}$$

Experiments on Grid'5000

588, 245 tetrahedra



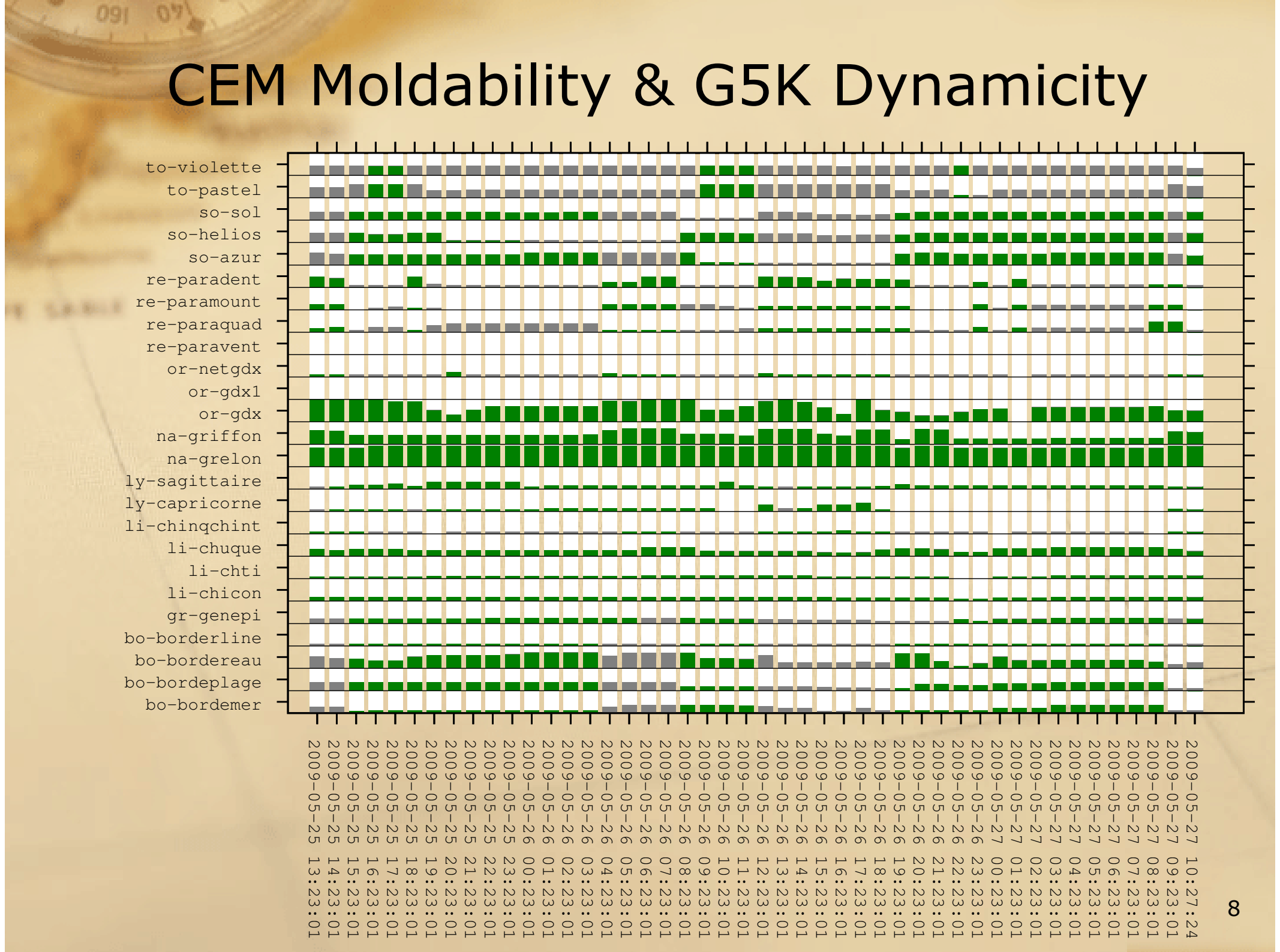
9, 659, 894 tetrahedra



- Expected optimal solution (all algorithms returned the same)
 - Added fast clusters
 - Removed slow clusters
- Actual optimal deployment slightly different

CEM Moldability & G5K Dynamicity

Task	Start Date	End Date	End Time
to-violette	2009-05-25	2009-05-27	10:27:24
to-pastel	2009-05-25	2009-05-27	09:23:01
so-sol	2009-05-25	2009-05-27	08:23:01
so-helios	2009-05-25	2009-05-27	07:23:01
so-azur	2009-05-25	2009-05-27	06:23:01
re-paradent	2009-05-25	2009-05-27	05:23:01
re-paramount	2009-05-25	2009-05-27	04:23:01
re-paraquad	2009-05-25	2009-05-27	03:23:01
re-paravent	2009-05-25	2009-05-27	02:23:01
or-netgdx	2009-05-25	2009-05-27	01:23:01
or-gdx1	2009-05-25	2009-05-27	00:23:01
or-gdx	2009-05-25	2009-05-27	23:23:01
na-griffon	2009-05-25	2009-05-27	22:23:01
na-grelon	2009-05-25	2009-05-27	21:23:01
ly-sagittaire	2009-05-25	2009-05-27	20:23:01
ly-capricorne	2009-05-25	2009-05-27	19:23:01
li-chinqchint	2009-05-25	2009-05-27	18:23:01
li-chuque	2009-05-25	2009-05-27	17:23:01
li-cti	2009-05-25	2009-05-27	16:23:01
li-chicon	2009-05-25	2009-05-27	15:23:01
gr-genepi	2009-05-25	2009-05-27	14:23:01
bo-borderline	2009-05-25	2009-05-27	13:23:01
bo-bordereau	2009-05-25	2009-05-27	12:23:01
bo-bordeplage	2009-05-25	2009-05-27	11:23:01
bo-bordemer	2009-05-25	2009-05-27	10:23:01



Problem Statement

- ❑ How to automatically select the best resources for executing a moldable application given that the objective function to minimize is run-specific
 - ❑ Minimize makespan
 - ❑ Minimize cost
 - ❑ Minimize energy
 - ❑ Finish before a given deadline
 - ❑ Etc..
- ❑ And that
 - ❑ resources are efficiently used and
 - ❑ the availability of resources is guaranteed

Limitations of Existing RMS

❑ Traditional RMS (Cluster, Supercomputers, ...)

- ❑ 1 queue $\Leftrightarrow$ 1 known set of machines
 - ❑ **Fixed sized jobs**
 - ❑ Eg. 1024 nodes of quad-core processors
 - ❑ **Rigid jobs (JSDL)**
 - ❑ User selected number of processor
 - ❑ **Moldable jobs (SLURM, OAR, ...)**
 - ❑ User selected range of number of processor
 - ❑ 1 global walltime (e.g., SLURM)
 - ❑ A set of rigid jobs (e.g., OAR)



❑ Federation of clusters / Grids

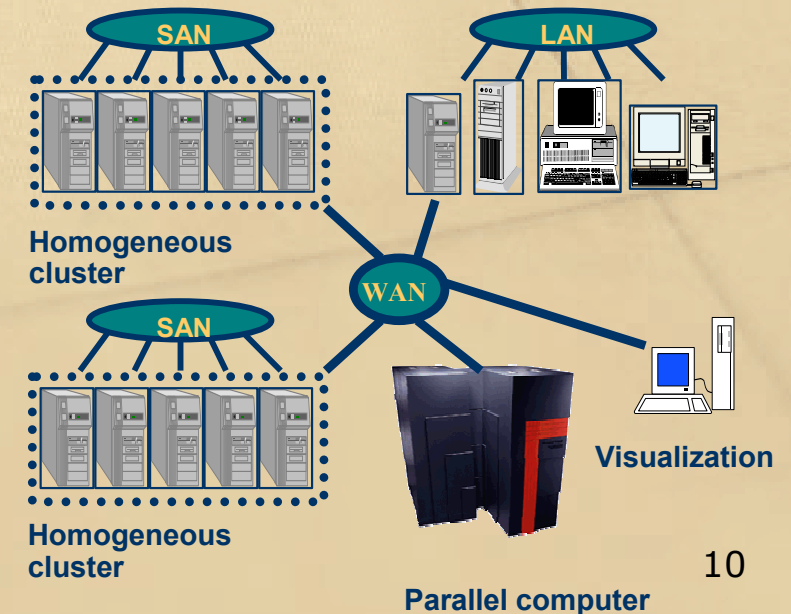
- ❑ Finding adequate resource
 - ❑ Constraint or match-making language

❑ Clouds

- ❑ On demand VMs
 - ❑ No guarantee on the availability

❑ Externally building a model of RMS

❑ Multiple submissions



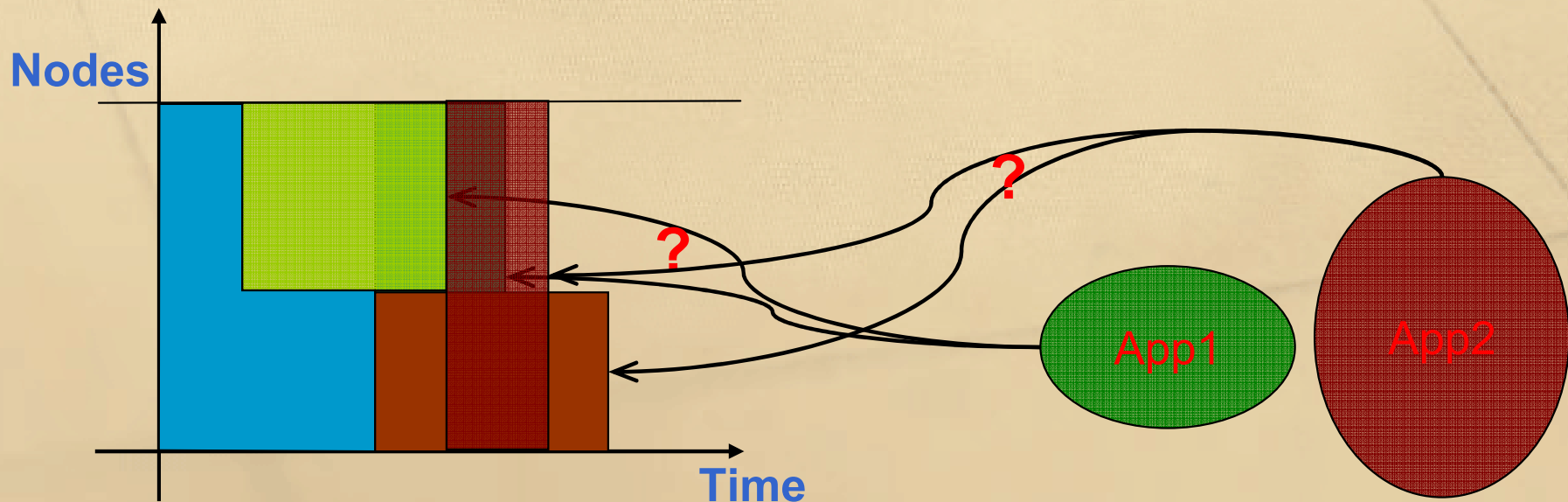
The background of the slide features a vintage-style map with a compass rose in the upper left corner. The map is aged and yellowed, with various geographical features and labels. The compass rose is circular with a needle pointing towards the top. The text "CooRMv1" is prominently displayed in the center-right of the slide.

CooRMv1

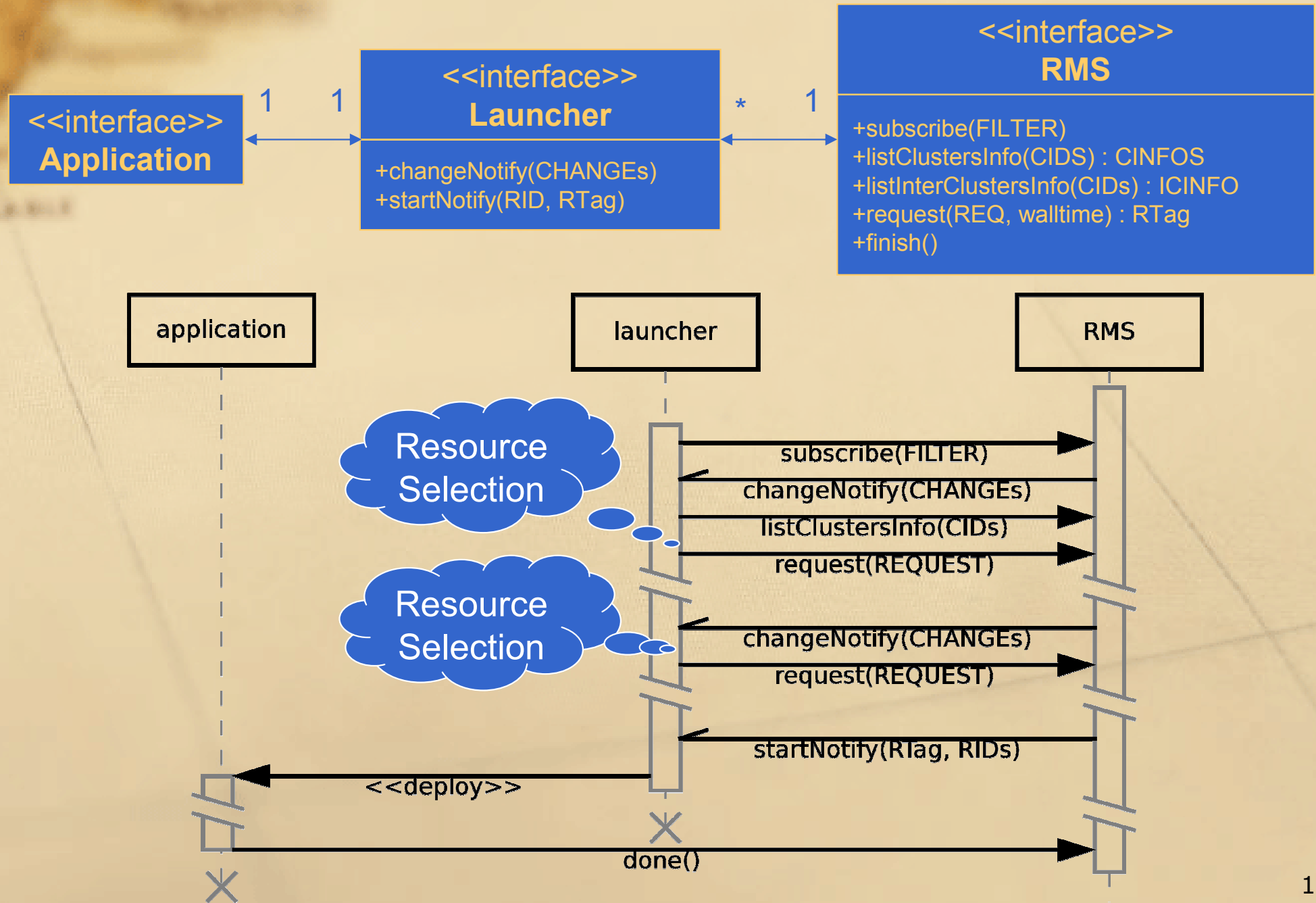
Supporting Complex Moldable
Applications

Untying RMS from Moldable Application Scheduling

- ❑ Applications should take a more active role in the scheduling
- ❑ RMS gives application the resource occupation (a view)
- Application decides what resources to request

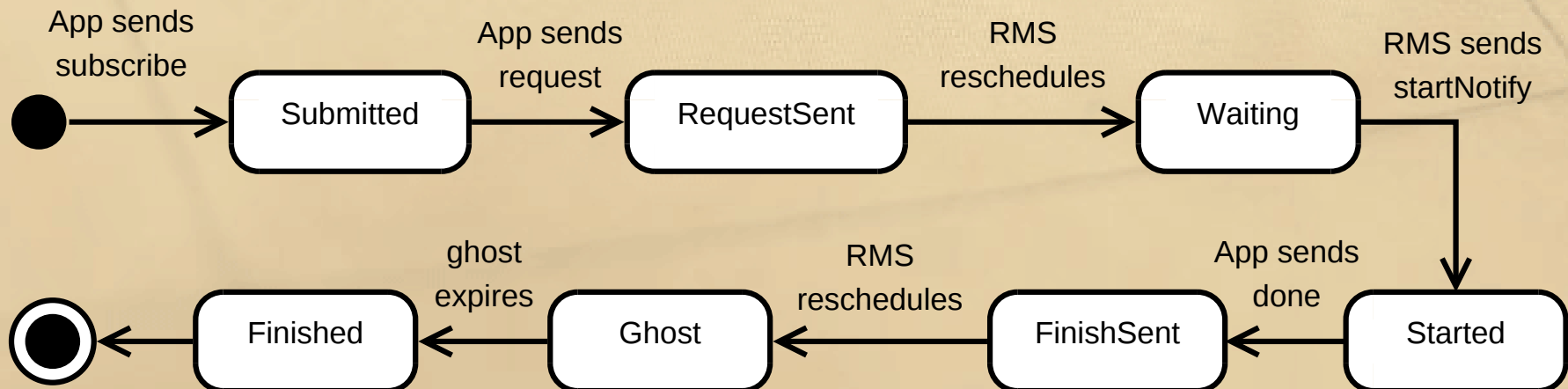
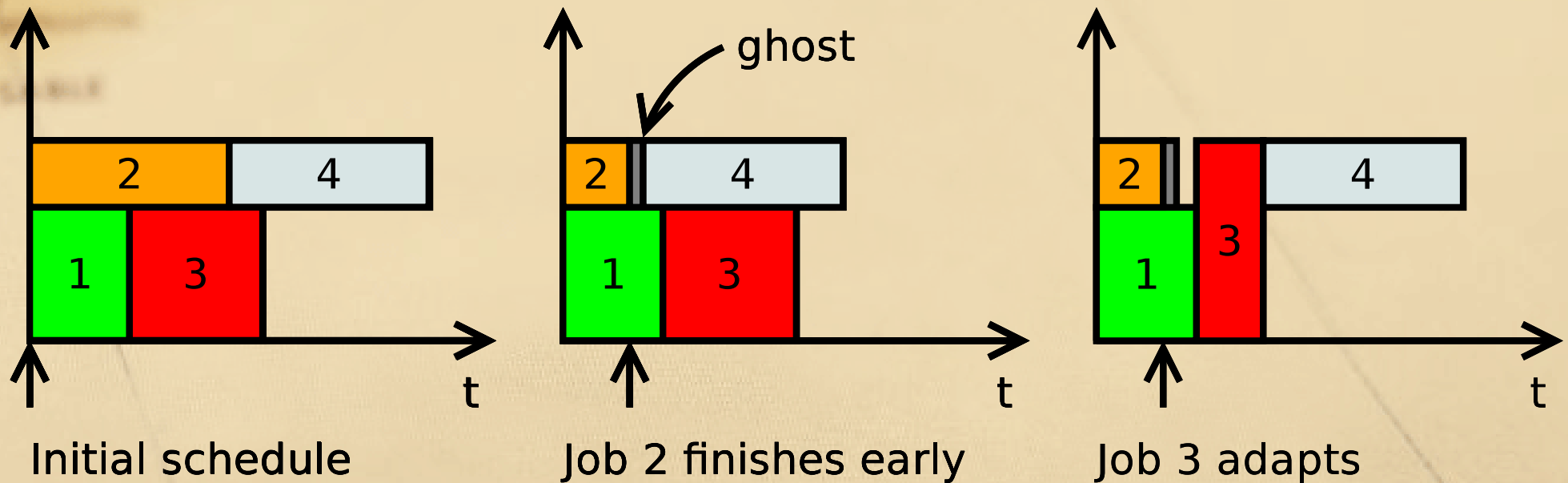


CooRMv1: Overview



RMS-side Scheduling

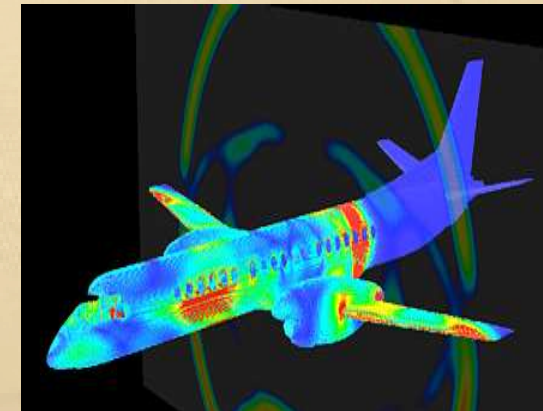
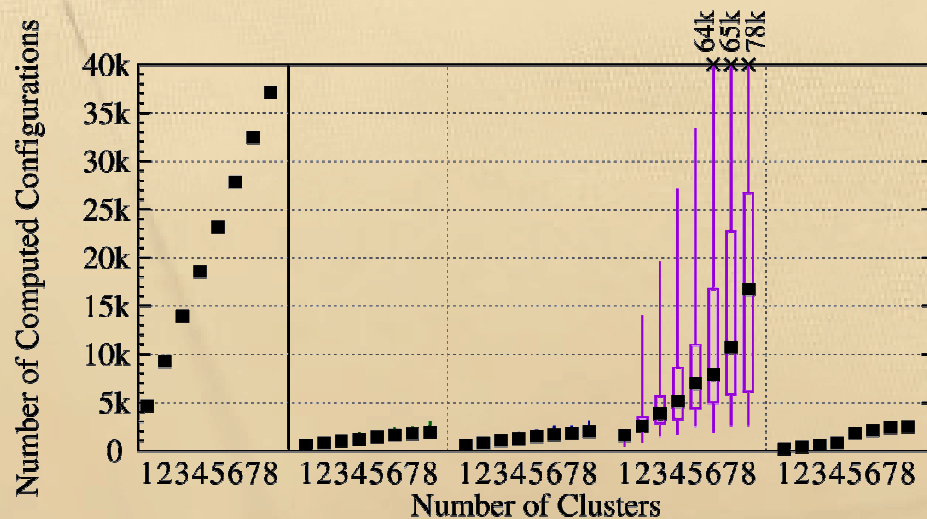
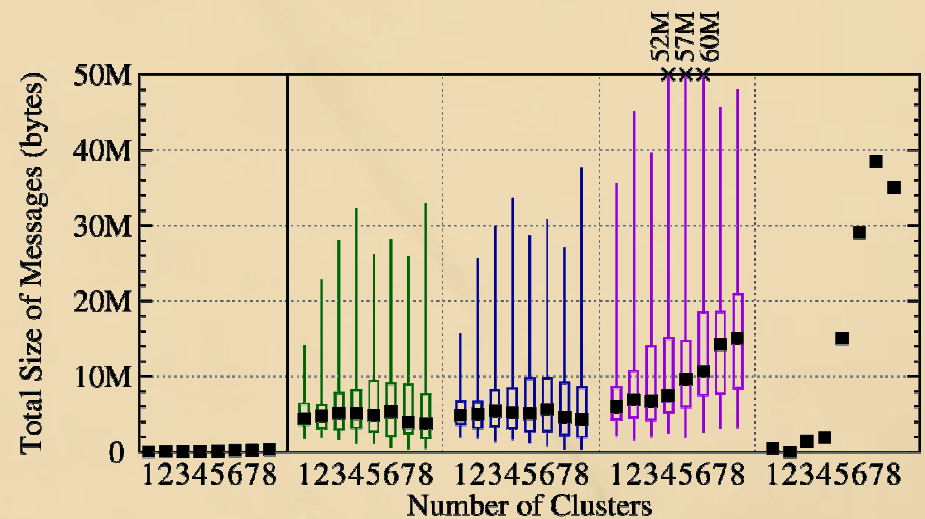
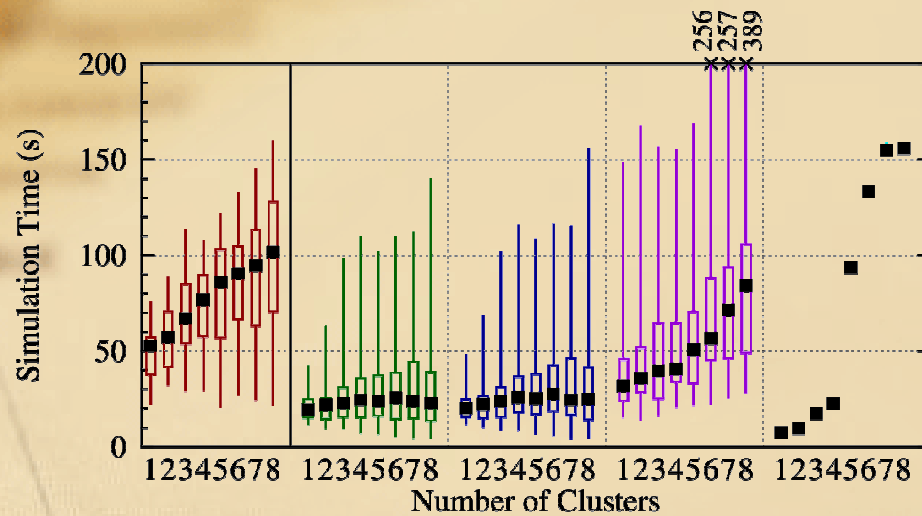
Fair-start Delay and Ghost Zones



RMS Scheduling Algorithm

1. Mark as **Ghost** applications which sent **Done**
2. Mark as **Finished** ghosts which have expired
3. Generate the initial last view by adding the **Started** applications and **Ghosts** to it.
4. For all Request Sent and Waiting applications, taken in the order of their submission time:
 1. Set the application's view to the last view;
 2. Compute the application's estimated start time;
 3. Add its request to the last view;
 4. Mark the application as Waiting.
5. Send **changeNotify** messages to applications whose view has changed because of the previous steps.

CooRM – Scalability

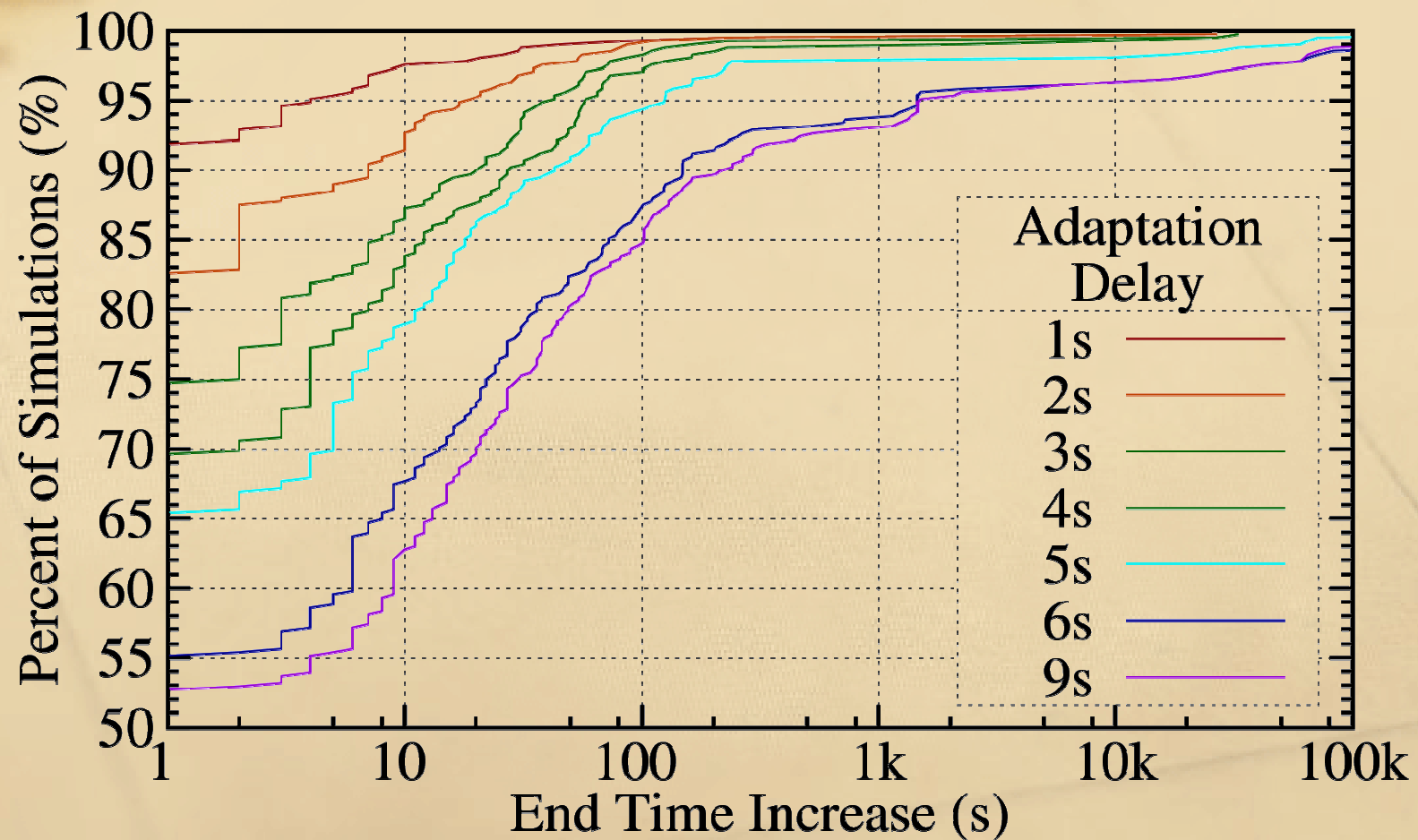


Legend

OAR : 0% CEM

CooRM: 0% CEM, 0.5% CEM, 50% CEM, 100% CEM

CooRM – Fairness



Content

- *Context*

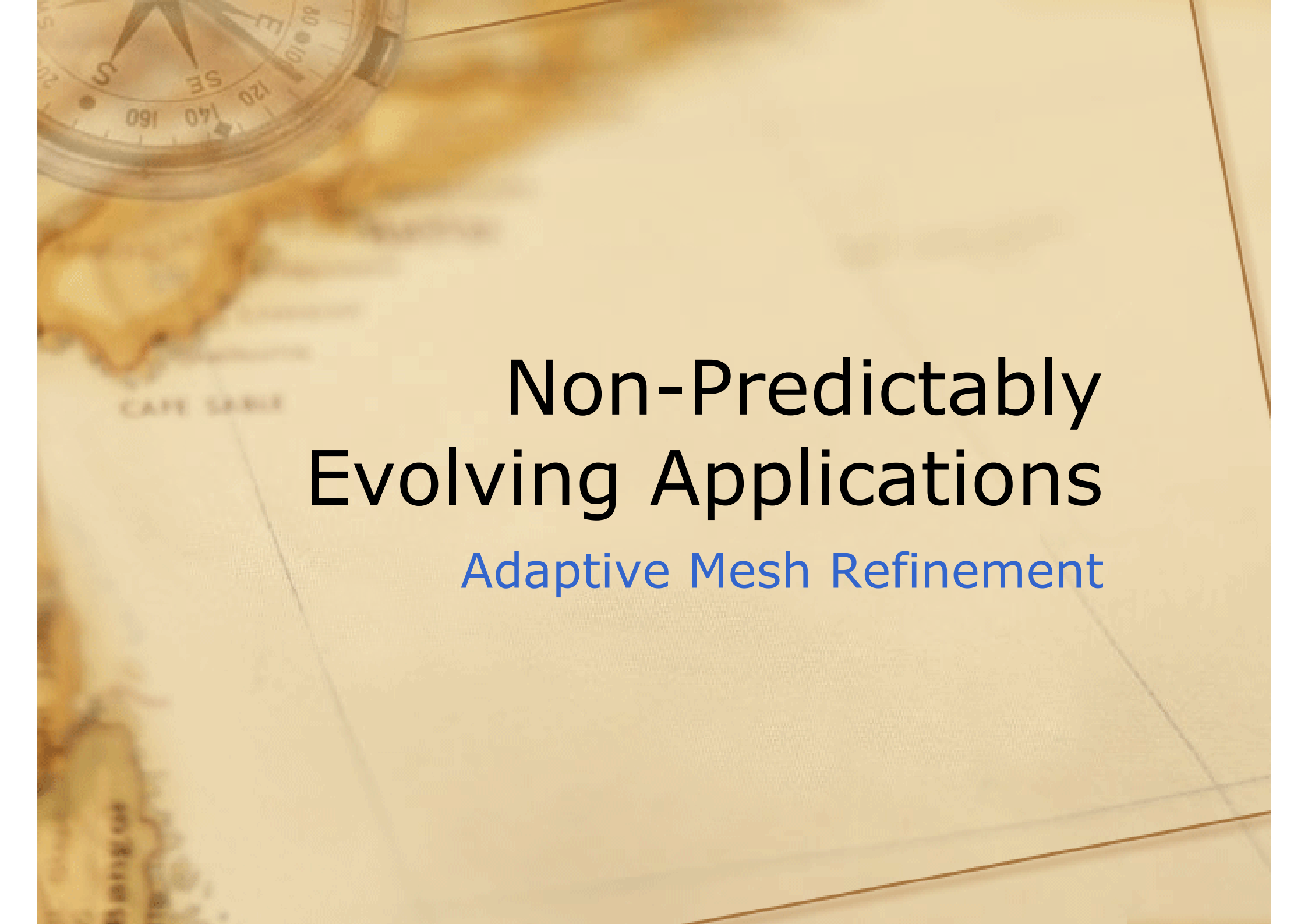
- *Complex moldable applications*

- *Example (Computational Electromagnetic)*
- *Limitations of existing RMS*
- *CooRMv1 – Model & Experiments*

- *Non predictably evolving applications*

- *Example (AMR)*
- *Limitations of existing RMS*
- *CooRMv2 – Model & Experiments*

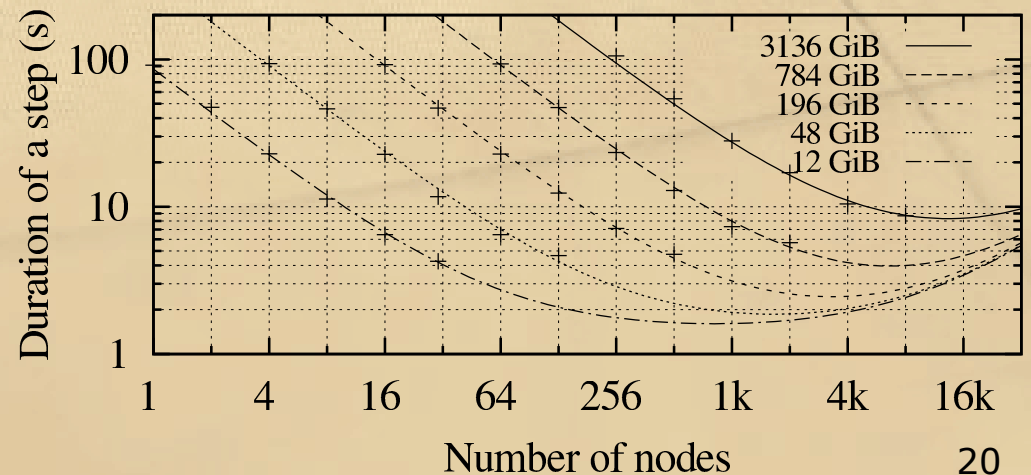
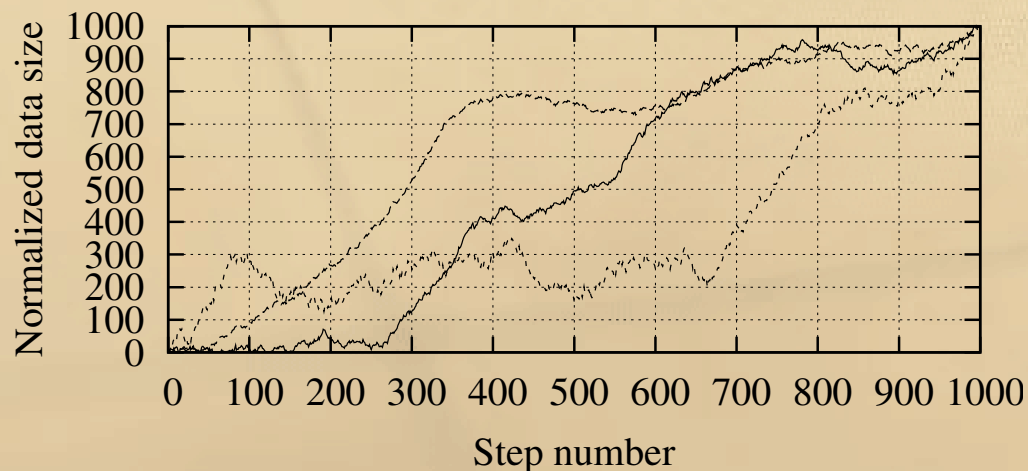
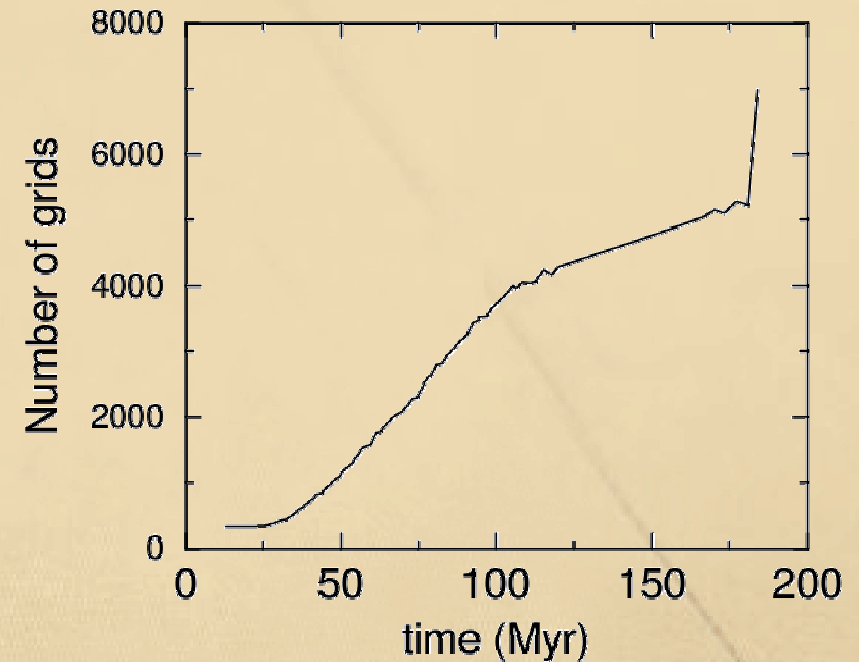
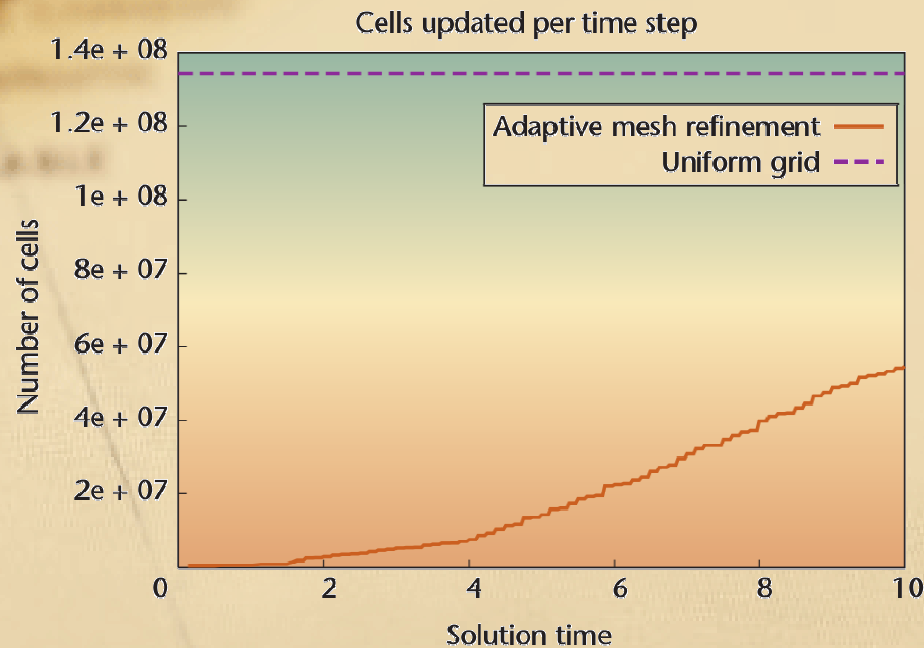
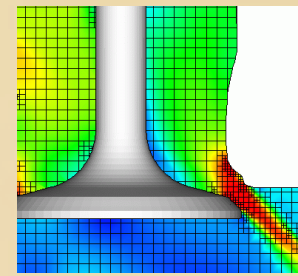
- *Conclusion and future works*

The background of the slide features a vintage-style map with a compass rose in the upper left corner. The map is aged and yellowed, with various geographical features and labels visible, though some are blurred. The compass rose shows cardinal and intercardinal directions. The overall aesthetic is that of an old, weathered document.

Non-Predictably Evolving Applications

Adaptive Mesh Refinement

AMR Applications



Problem Statement Revisited

- ❑ How to automatically select the best resources for executing a moldable **and evolving** application given that the objective function to minimize is run-specific
 - ❑ Minimize makespan
 - ❑ Minimize cost
 - ❑ Minimize energy
 - ❑ Finish before a given deadline
 - ❑ Etc..
- ❑ And that
 - ❑ resources are efficiently used and
 - ❑ the availability of resources is guaranteed

Limitations of Current RMS

- ❑ Malleable jobs

- ❑ RMS tells applications to grow/shrink

- ❑ Clouds

- ❑ The illusion of infinite computing resources available on demand"
 - ❑ Infinite? Actually up to 20
 - ❑ Even without this limit: "Out of capacity" errors
 - ❑ Application may run out-of-memory

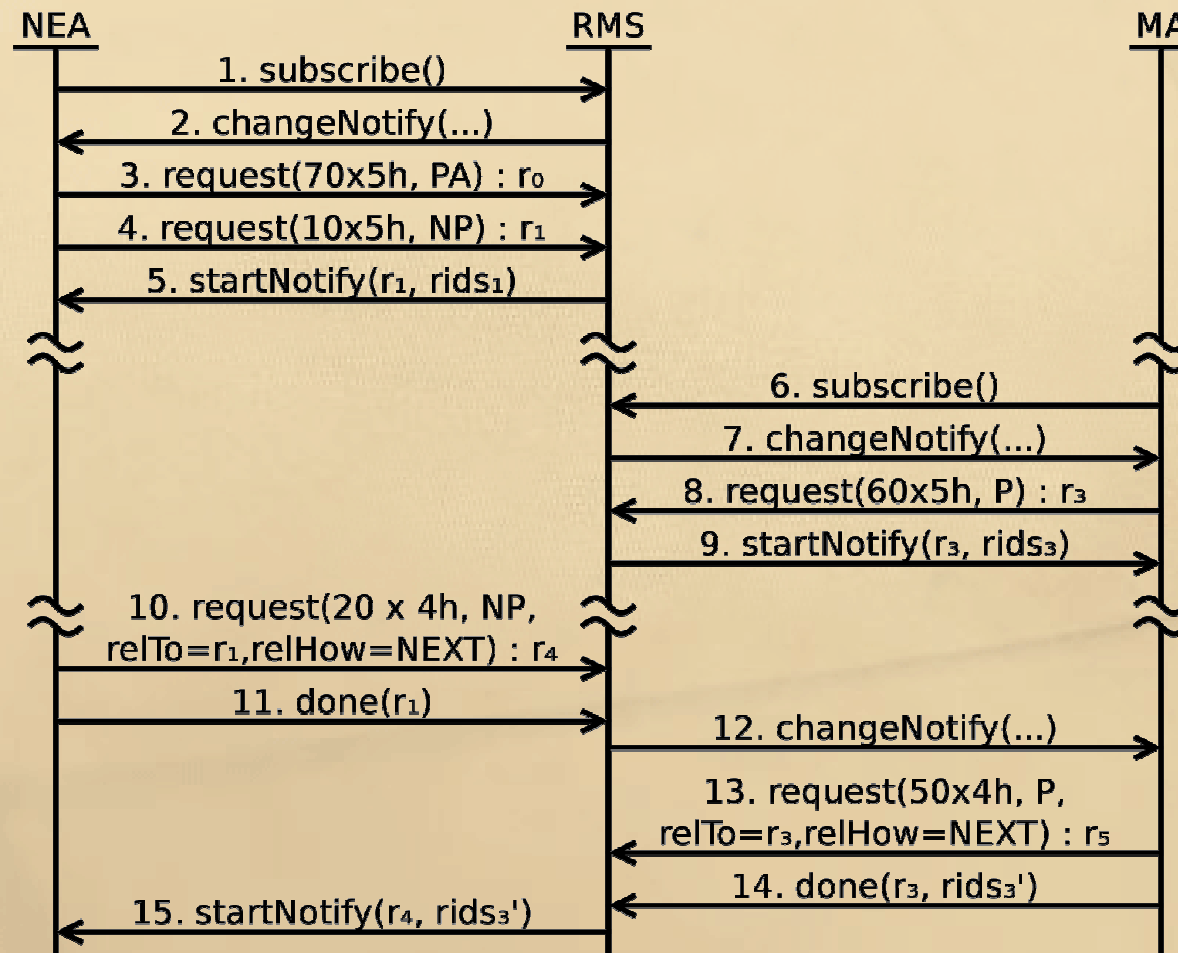
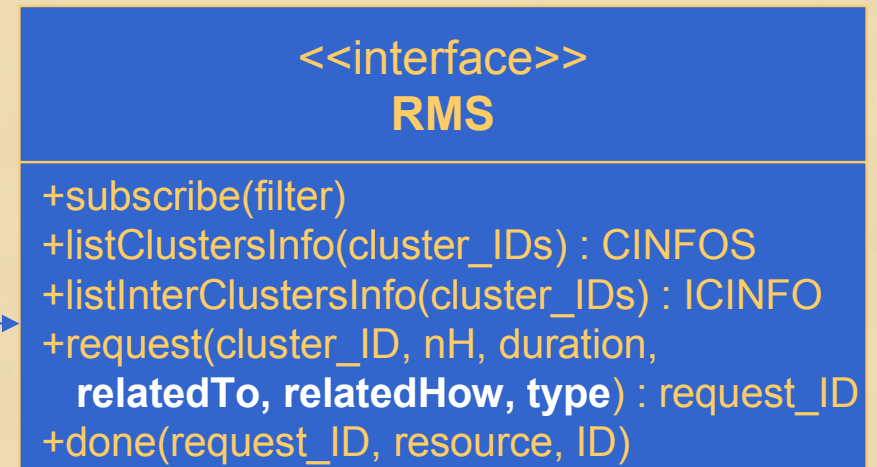
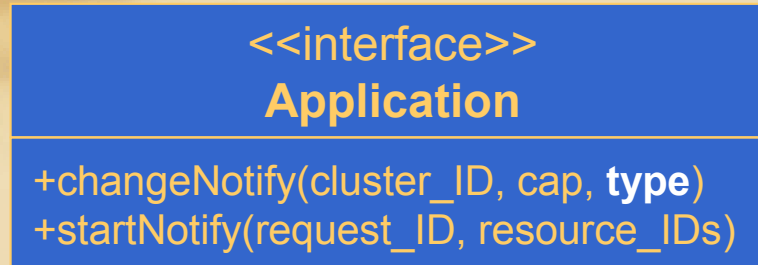
- ❑ Ideally, RMS guarantees the availability of resources to an AMR application?



CooRMv2

Supporting Non-Predictably
Evolving Applications

CooRMv2



Resource Request Model

- ❑ CooRMv1 extension
- ❑ New request parameters
 - ❑ Cluster ID, number of nodes, duration, **type**
 - ❑ RMS chooses start time
 - ❑ Node IDs are allocated to the application
- ❑ Type
 - ❑ Non-preemptible (default in major RMSs)
 - ❑ Preemptible (think OAR best-effort jobs)
 - ❑ Pre-allocation
 - ❑ “I do not currently need these resources, but make sure I can get them immediately if I need them.”

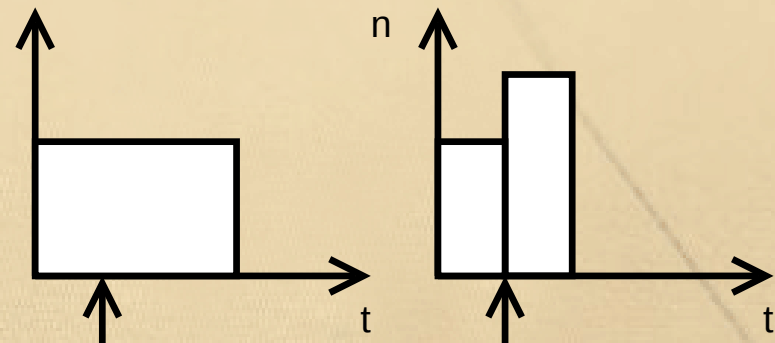
High Level Operations

- Low level operations

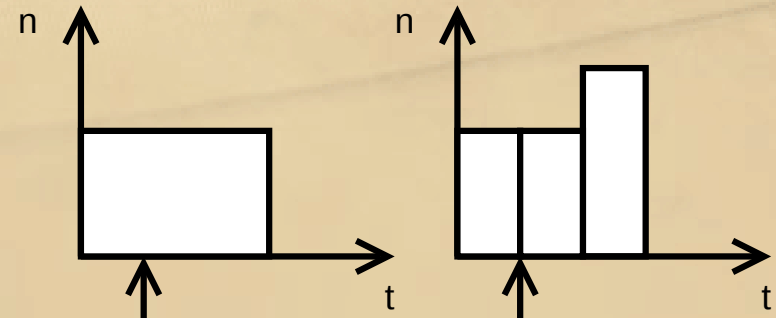
 - Not very expressive but simplify scheduling

- High level operations

 - Spontaneous update*



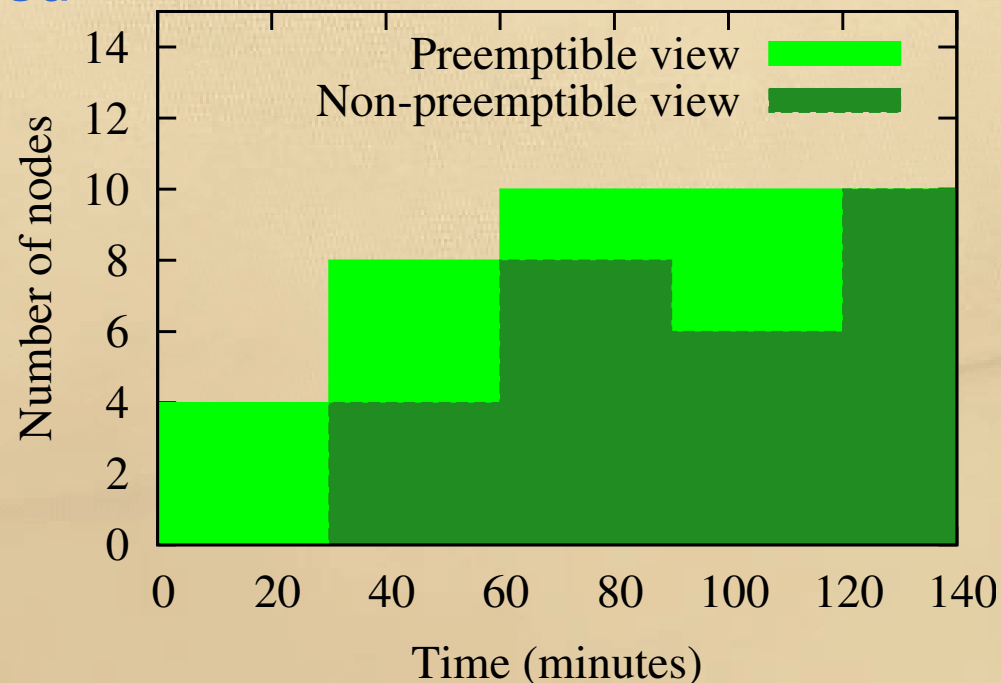
 - Announced updates*



*: Update is only guaranteed within a pre-allocation request. ²⁶

Views

- ❑ Apps need to adapt their requests to the availability of the resources
- ❑ Each app is presented with two views
 - ❑ non-preemptible
 - ❑ preemptible
- ❑ Preemptible view informs when resources need to be preempted

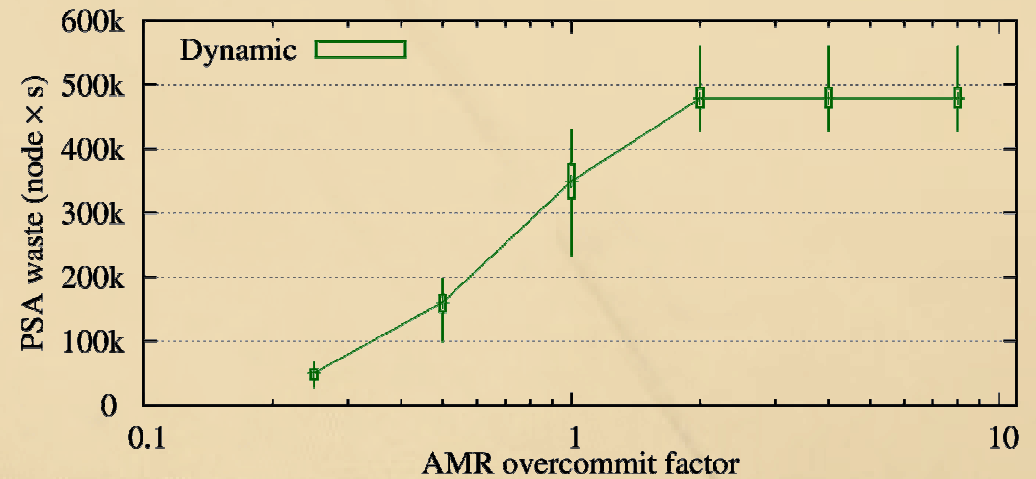
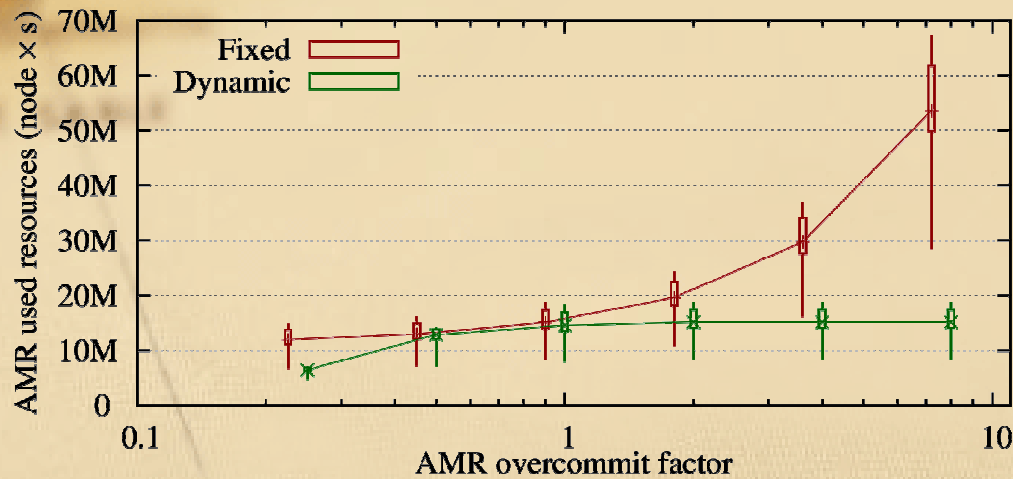


Scheduling Algorithm

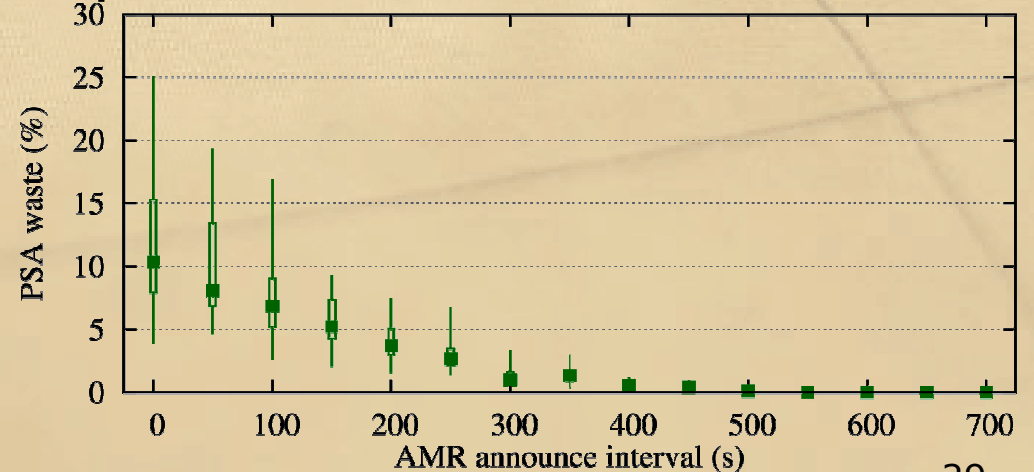
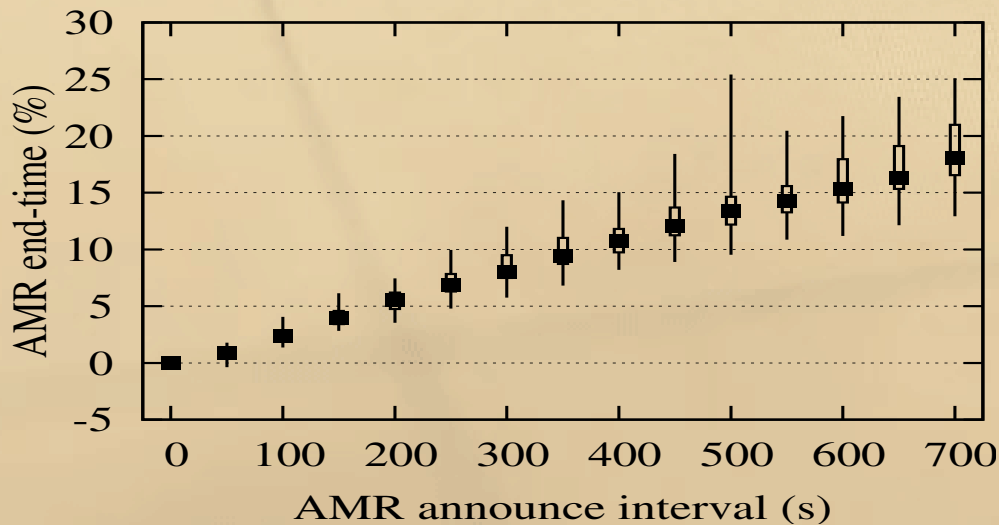
- ❑ Pre-allocations and non-preemptible requests
 - ❑ Conservative Back-Filling (CBF)
- ❑ Preemptible requests
 - ❑ Equi-partitioning

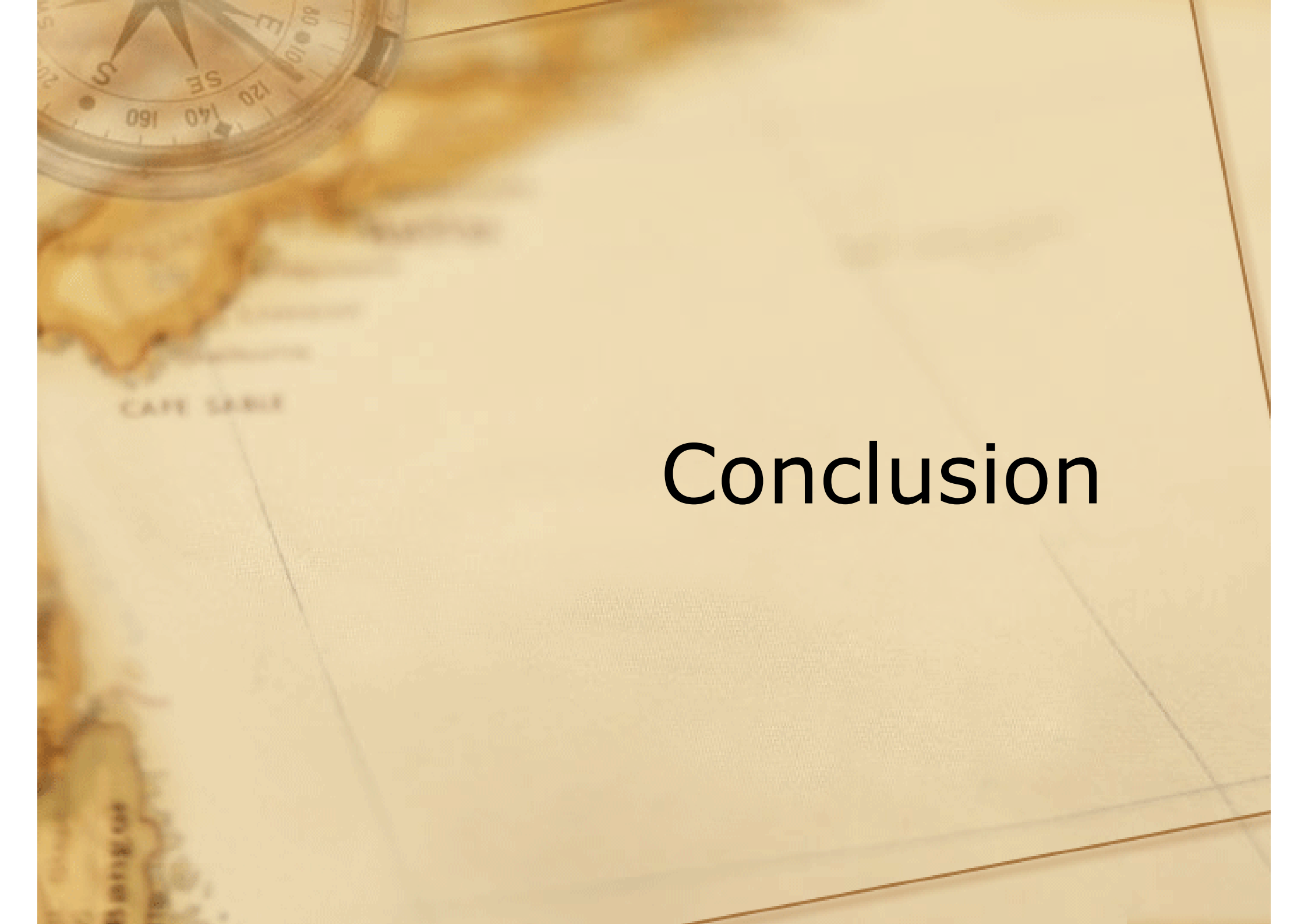
CooRMv2: AMR & PSA

Spontaneous Updates



Announced Updates





Conclusion

CooRMv1/2

❑ Simple & Generic RMS–Application Interface

- ❑ Application specific resource selection algorithm
 - ❑ Application specific optimization (time, cost, ...)
- ❑ Applications with dynamic resource requirements
 - ❑ Malleable and evolving
- ❑ Proof of concept implementation
 - ❑ Good performance

❑ Future Work

- ❑ Add (intra-cluster) network topology support
- ❑ Develop (generic) resource selection algorithms for more applications
- ❑ Integrate CooRM concept into existing RMS
 - ❑ OAR is planned
- ❑ Is CooRM suitable for distributed systems?
 - ❑ DIET? XtremOS?
- ❑ Is CooRM suitable for Clouds?