

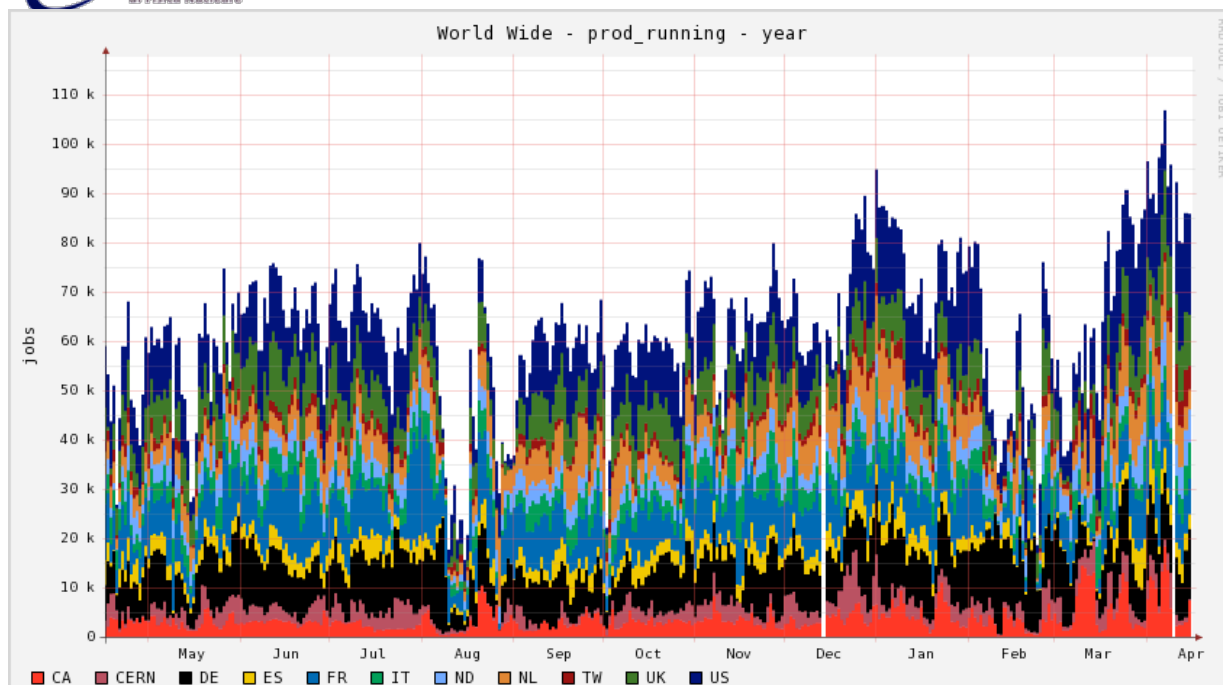
Stato del Calcolo e “produzioni” MC Italiane

Gianpaolo Carlino
INFN Napoli

Roma Tor Vergata, 17 Aprile 2012

- Attività di Computing in Atlas
- La distribuzione dei dati
- L'accesso ai dati
- Utilizzo risorse Italiane

Attività di Computing in ATLAS



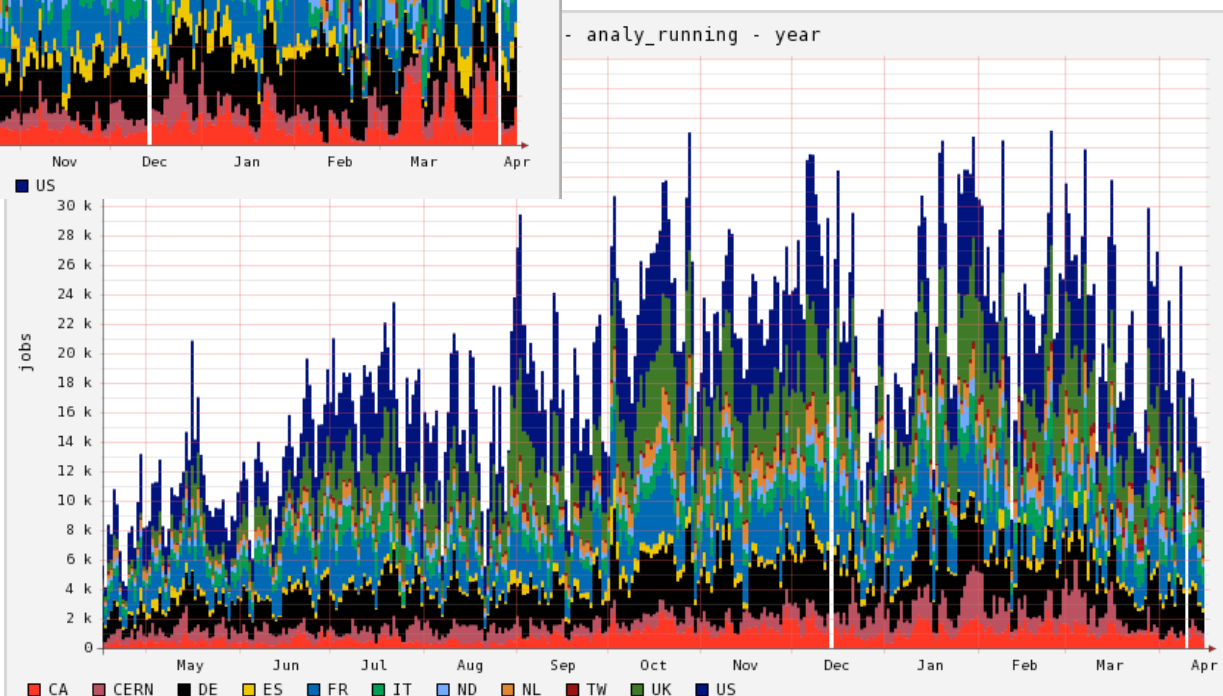
Numero di job simultanei nell'ultimo anno

Produzione:

- > 60k job, costante
- incremento inizio 2012 per reco MC11

Analisi:

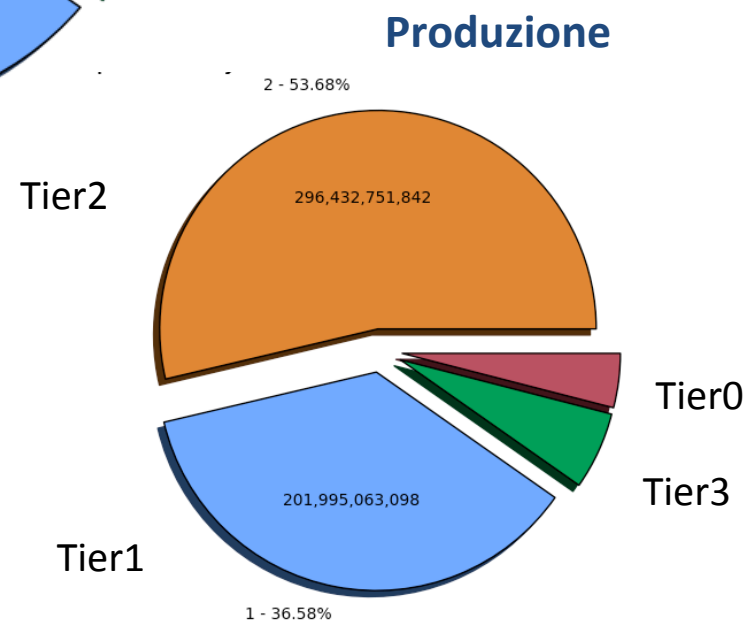
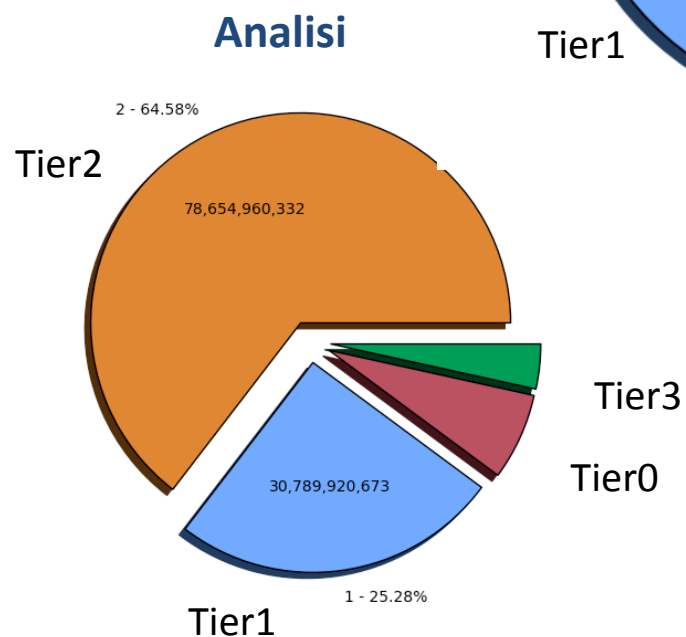
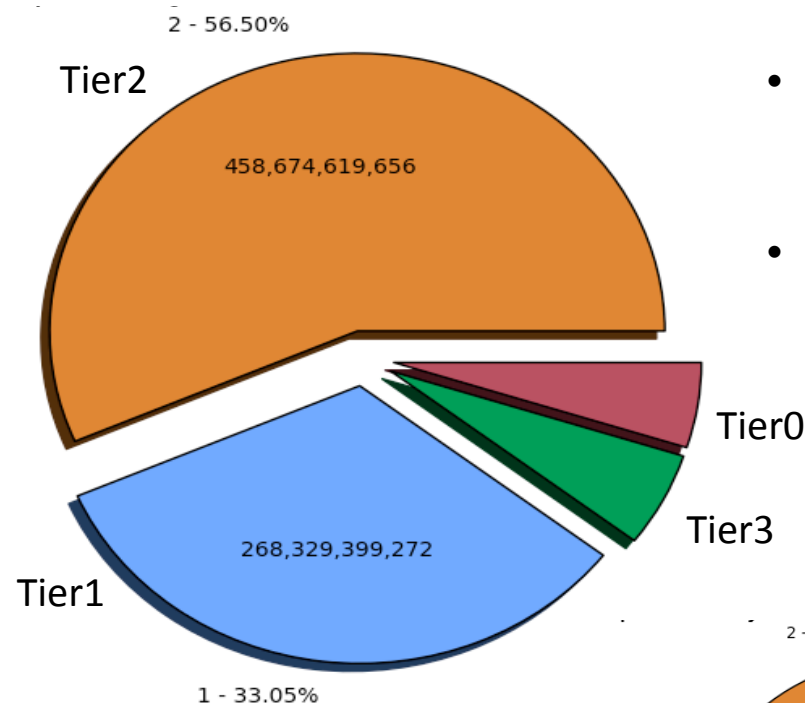
- crescita lineare nel 2011 fino a > 20k job
- Alta attività nel 2012 per le conferenze invernali
- In attesa di statistica per riprendere l'attività seria





Attività nei Tier

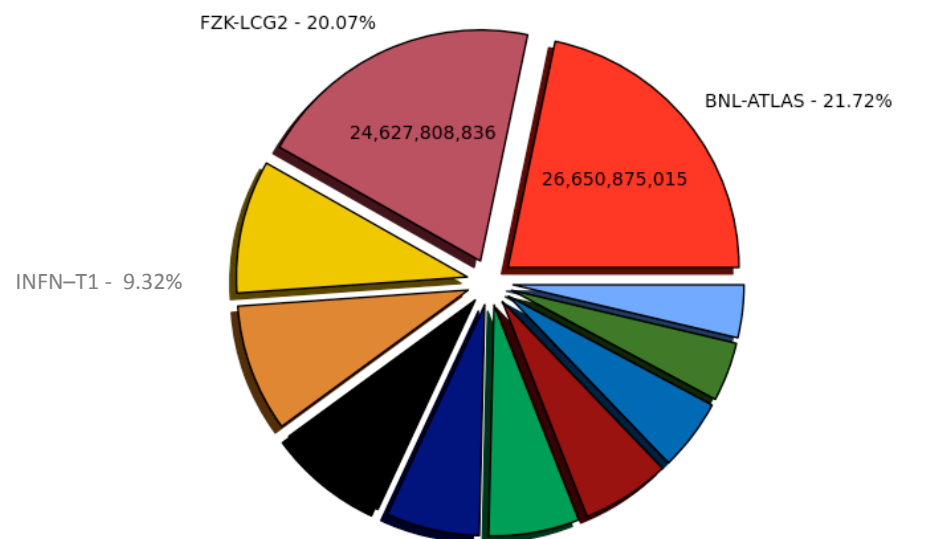
- I Tier2 forniscono la maggioranza delle risorse
- Contributo dei Tier3 non trascurabile



Utilizzo risorse al CNAF



CPU consumption All Jobs in seconds (Pie Graph) (Sum: 122,689,309,841)



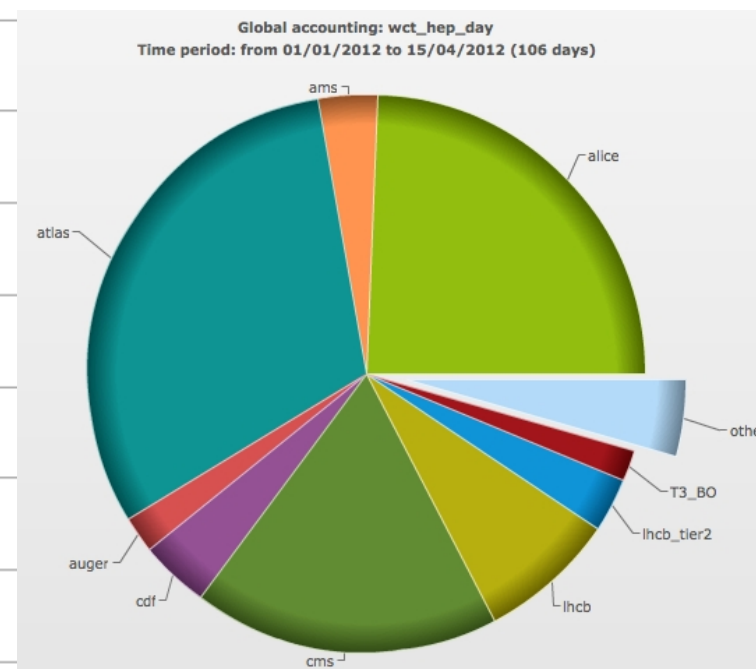
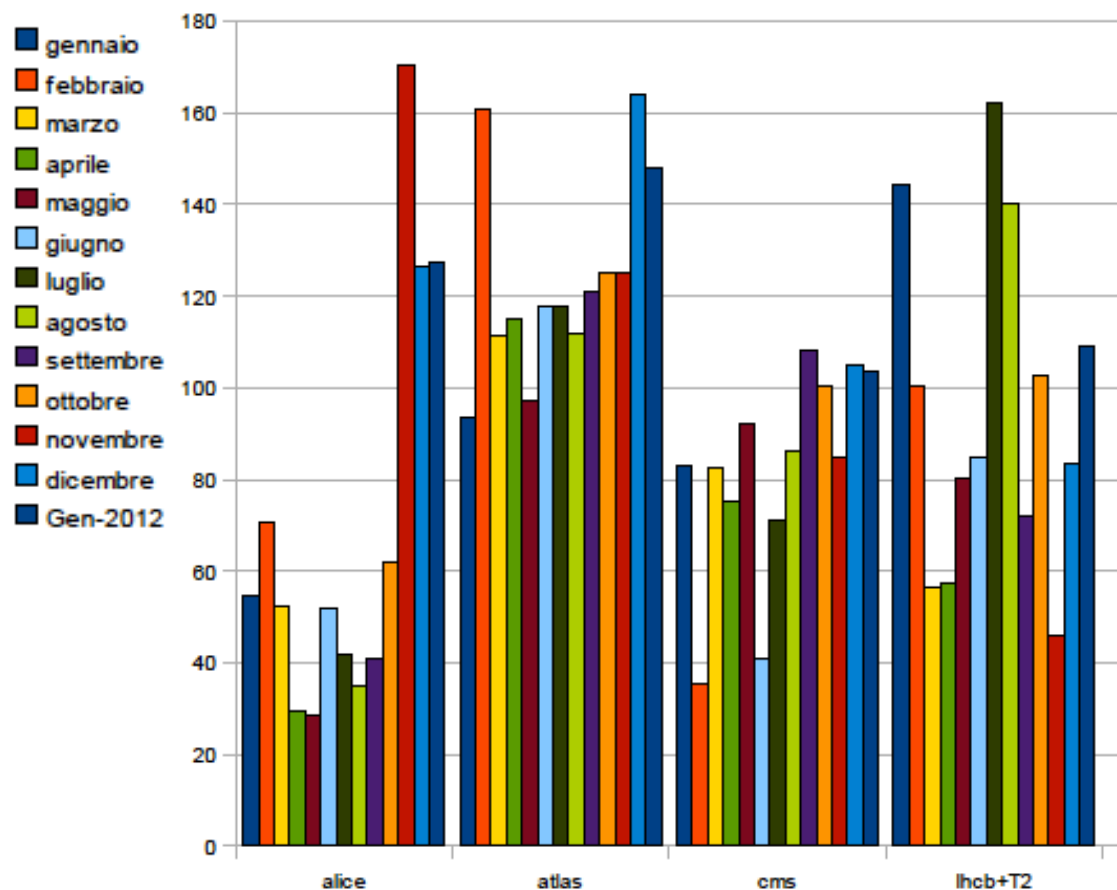
Il CNAF è tra i migliori Tier1 di ATLAS e lotta per la terza posizione

BNL-ATLAS - 21.72% (26,650,875,016)	FZK-LCG2 - 20.07% (24,627,808,836)	INFN-T1 - 9.32% (11,428,980,229)
RAL-LCG2 - 8.94% (10,965,027,754)	IN2P3-CC - 8.02% (9,843,646,413)	TRIUMF-LCG2 - 6.56% (8,047,972,281)
NIKHEF-ELPROD - 6.34% (7,776,041,913)	SARA-MATRIX - 6.20% (7,609,294,913)	NDGF-T1 - 5.00% (6,138,961,213)
TAIWAN-LCG2 - 4.13% (5,071,873,825)	PIC - 3.69% (4,528,827,448)	

Utilizzo risorse al CNAF

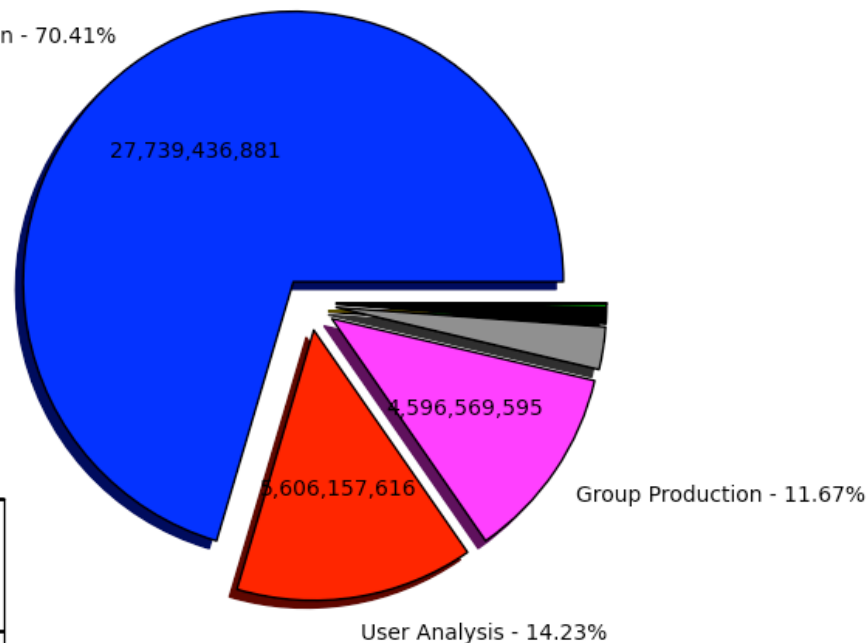
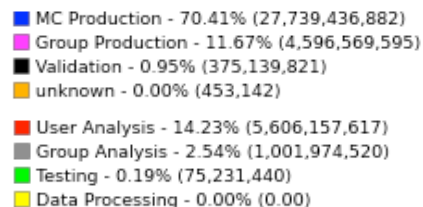


WCT HEP-SPEC day medi / Pledged 2011



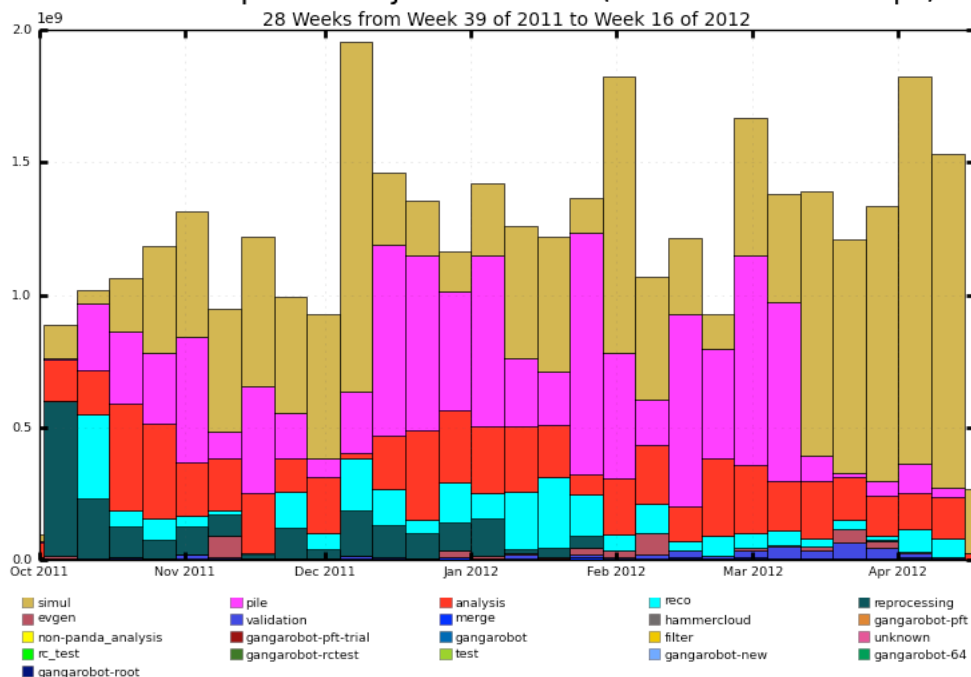


Uso delle risorse suddiviso per attività



simile agli altri Tier1

CPU consumption Good Jobs in seconds (Time Stacked Bar Graph)

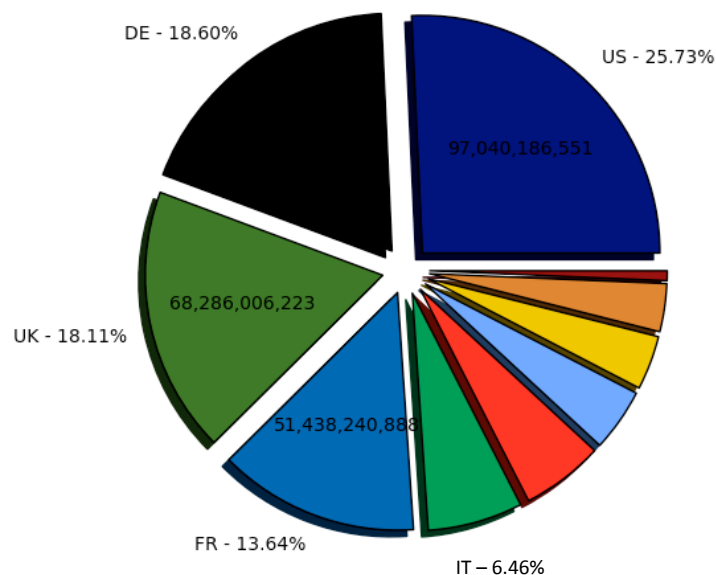


Corretto rapporto prod/analisi

Uso risorse nei Tier2

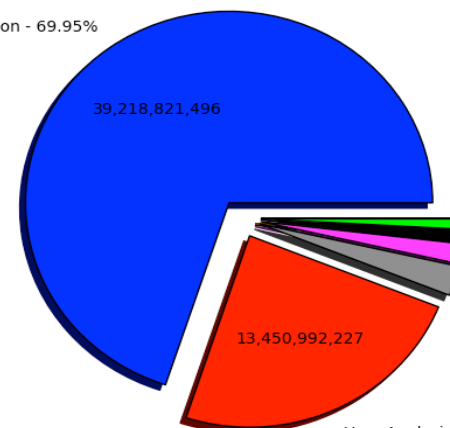


Uso risorse per "Processing Cloud"



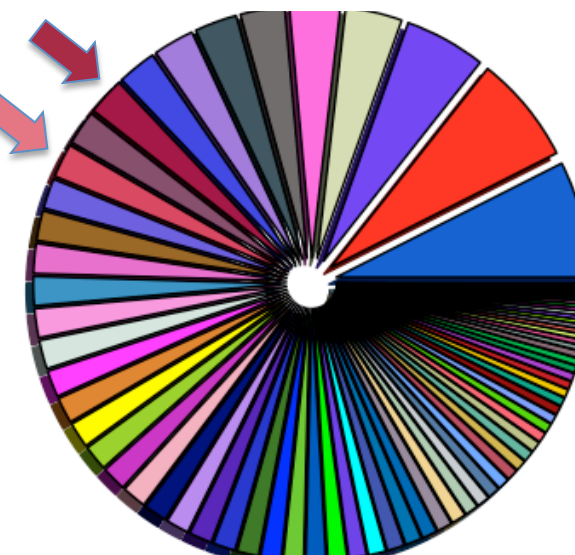
US - 25.73% (97,040,186,551)	DE - 18.60% (70,132,499,143)	UK - 18.11% (68,286,006,223)	FR - 13.64% (51,438,240,888)
IT - 6.46% (24,374,720,084)	CA - 5.59% (21,067,897,417)	ND - 4.28% (16,131,461,686)	ES - 3.68% (13,864,013,145)
NL - 3.26% (12,310,987,750)	TW - 0.67% (2,510,830,042)		

MC Production - 69.95%



per
attività

per Tier2

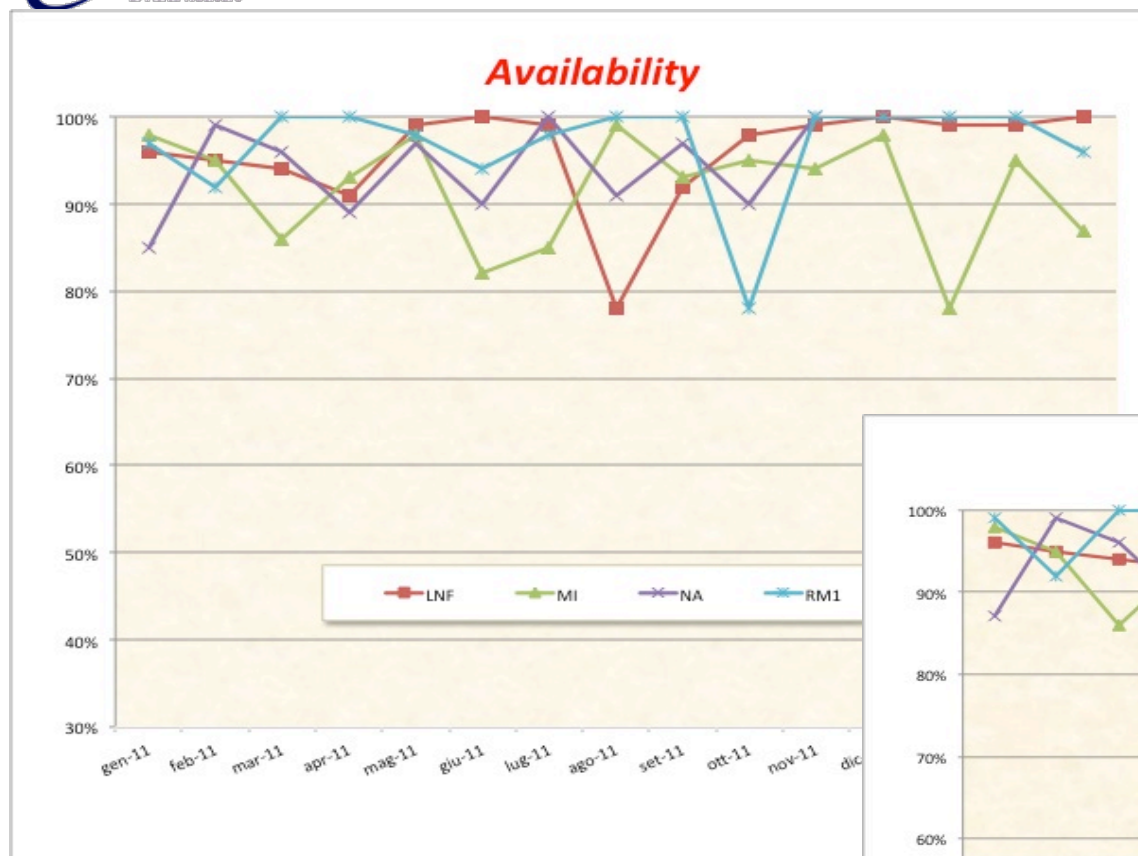


MWT2_UC - 7.45% (28,090,192,848)	AGLT2 - 6.93% (26,126,707,318)
WT2 - 5.20% (19,611,599,440)	UKI-LT2-QMUL - 3.62% (13,645,826,630)
UKI-NORTHGRID-MAN-HEP - 3.22% (12,160,150,435)	DESY-HH - 2.83% (10,667,339,532)
SIGNET - 2.71% (10,237,951,212)	LRZ-LMU - 2.61% (9,857,433,342)
IN2P3-CC-T2 - 2.49% (9,391,401,268)	INFN-NAPOLI-ATLAS - 2.39% (9,018,518,626)
MPPMU - 2.19% (8,272,888,360)	INFN-ROMA1 - 2.16% (8,158,771,754)
UKI-NORTHGRID-LANCS-HEP - 1.91% (7,213,388,680)	TOKYO-LCG2 - 1.91% (7,191,440,361)
UTA_SWT2 - 1.90% (7,158,135,924)	PRAGUE-LCG2 - 1.86% (6,996,313,121)
UKI-SCOTGRID-GLASGOW - 1.85% (6,974,327,307)	SWT2_CPB - 1.83% (6,898,743,849)
GRIF-IRFU - 1.68% (6,342,541,203)	plus 60 more

Quinta cloud

La percentuale può essere molto diversa dai pledges dichiarati a causa delle risorse a disposizione nelle varie per le attività nazionali (anche in IT)

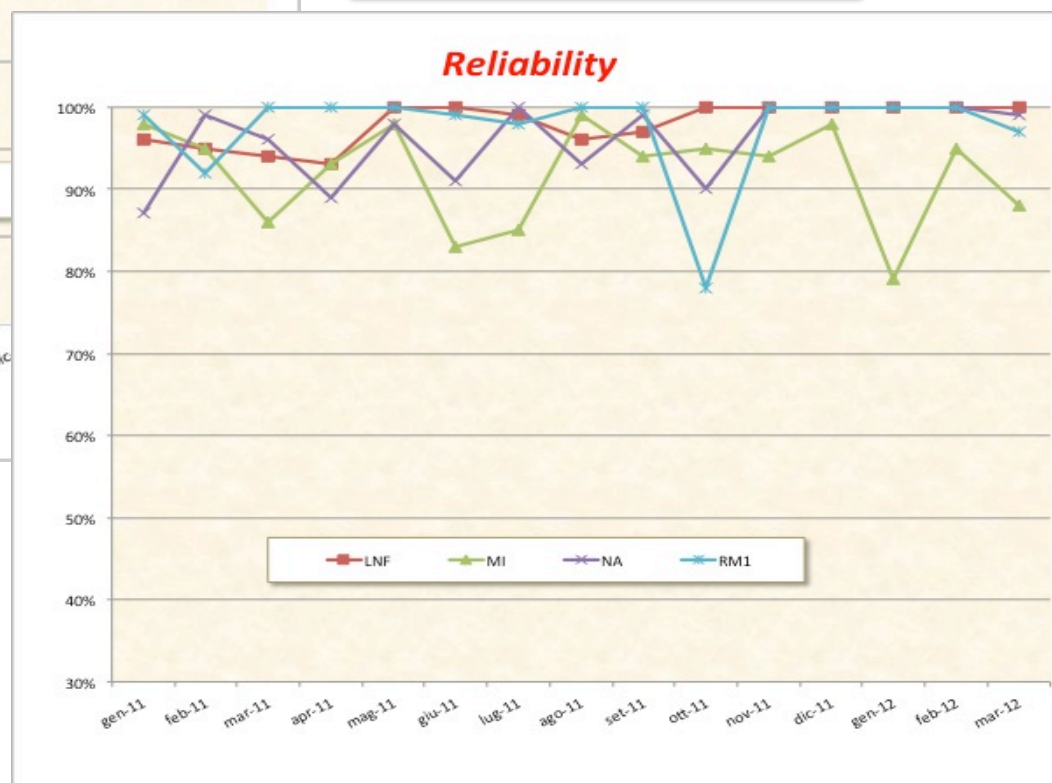
Reliability & Availability 2011-12



Availability =
time_site_is_available/total_time

Reliability =
time_site_is_available/
(total_time-time_site_is_sched_down)

Valori medi 2011-12			
Frascati		Milano	
rel	ava	rel	ava
98%	96%	92%	92%
Napoli		Roma	
rel	ava	rel	ava
96%	95%	98%	97%



Classificazione dei Tier2



- Necessità di individuare i siti più affidabili per l'analisi cui inviare la maggior parte dei dati.
- Classificazione in base alle performance (stabilità)

4 Gruppi

- Alpha: (60% share): T2D con rel > 90%
- Bravo: (30% share): non T2D con rel > 90%
- Charlie: (10% share): 80% < rel < 90%
- Delta: (0% share): rel < 80%

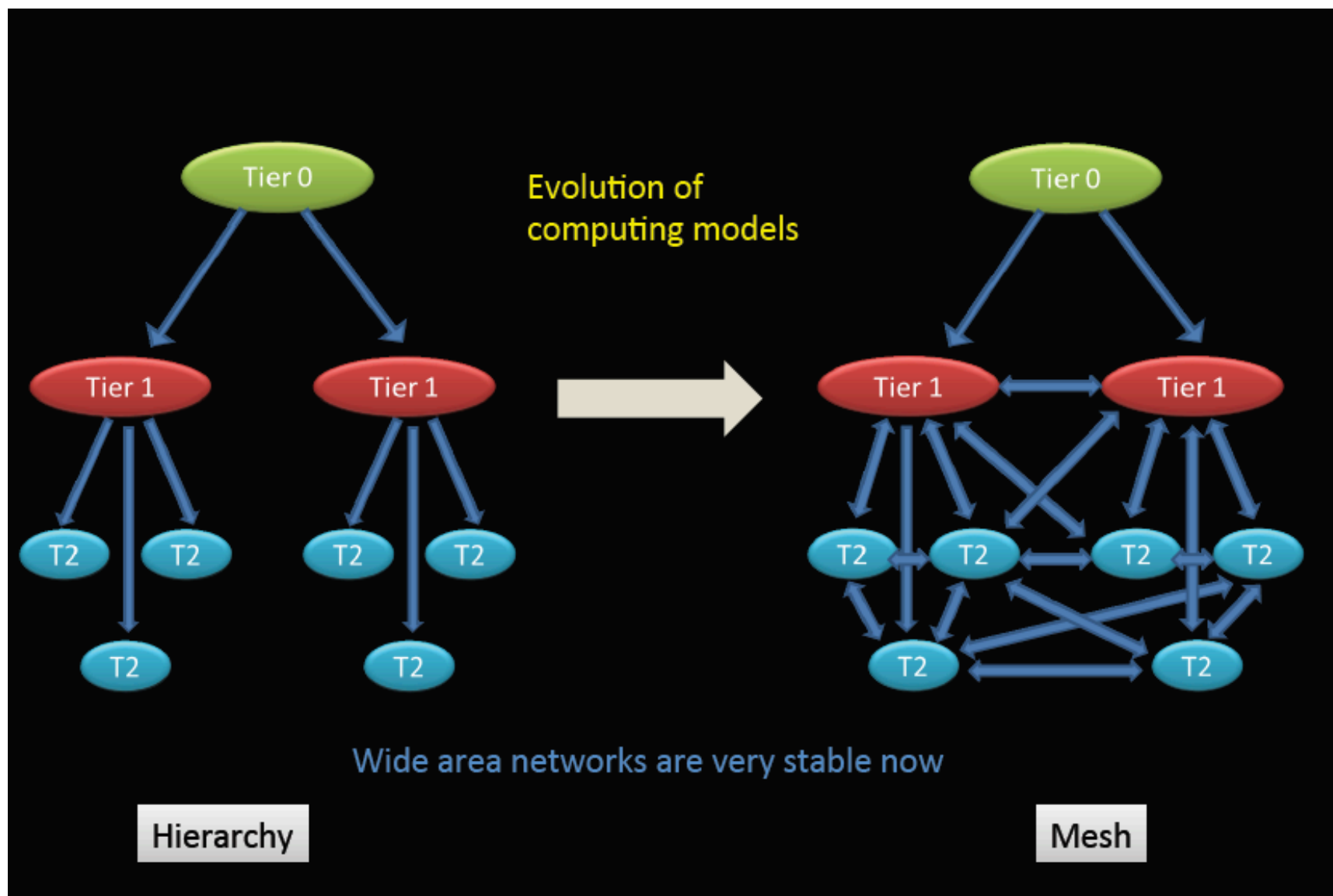
T1	tier	%	T2 Group	tiers status
CA	AUSTRALIA-ATLAS_DATADISK	1.2 %	bravo	registered: 6 ready: 6
	CA-ALBERTA-WESTGRID-T2_DATADISK	1.2 %	bravo	
	CA-MCGILL-CLUMEQ-T2_DATADISK	1.2 %	bravo	
	CA-SCINET-T2_DATADISK	2.4 %	alpha	
	CA-VICTORIA-WESTGRID-T2_DATADISK	2.4 %	alpha	
	SFU-LCG2_DATADISK	2.4 %	alpha	
DE	CSCS-LCG2_DATADISK	2.4 %	alpha	registered: 11 ready: 11
	CYFRONET-LCG2_DATADISK	0.6 %	charlie	
	DESY-HH_DATADISK	2.4 %	alpha	
	DESY-ZN_DATADISK	2.4 %	alpha	
	GOEGRID_DATADISK	2.4 %	alpha	
	HEPHY-UIBK_DATADISK	1.2 %	bravo	
	LRZ-LMU_DATADISK	2.4 %	alpha	
	MPPMU_DATADISK	2.4 %	alpha	
	PRAGUE-LCG2_DATADISK	1.2 %	bravo	
	UNI-FREIBURG_DATADISK	2.4 %	alpha	
ES	IFAE_DATADISK	2.4 %	alpha	registered: 6 ready: 6
	IFIC-LCG2_DATADISK	2.4 %	alpha	
	LIP-COIMBRA_DATADISK	0 %	delta	
	LIP-LISBON_DATADISK	2.4 %	alpha	
	NCG-INGRID-FT_DATADISK	2.4 %	alpha	
FR	UAM-LCG2_DATADISK	2.4 %	alpha	registered: 11 ready: 11
	BEIJING-LCG2_DATADISK	2.4 %	alpha	
	GRIF-IRFU_DATADISK	0 %	delta	
	GRIF-LAL_DATADISK	0.6 %	charlie	
	GRIF-LPHE_DATADISK	2.4 %	alpha	
	IN2P3-CPPM_DATADISK	1.2 %	bravo	
	IN2P3-LAPP_DATADISK	2.4 %	alpha	
	IN2P3-LPC_DATADISK	0 %	delta	
	IN2P3-EPSC_DATADISK	2.4 %	alpha	
	RO-02-NIPNE_DATADISK	0 %	delta	
	RO-07-NIPNE_DATADISK	0 %	delta	
	TOKYO-LCG2_DATADISK	1.2 %	bravo	

IT	INFN-FRASCATI_DATADISK	1.2 %	bravo	registered: 4 ready: 4
	INFN-MILANO-ATLAS_DATADISK	0.6 %	alpha	
	INFN-NAPOLI-ATLAS_DATADISK	2.4 %	alpha	
	INFN-ROMA1_DATADISK	2.4 %	alpha	
NL	RRC-KI_DATADISK	0 %	delta	registered: 9 ready: 8
	CSTCDIE_DATADISK	0 %	delta	
	IL-TAU-HEP_DATADISK	1.2 %	bravo	
	JINR-LCG2_DATADISK	1.2 %	bravo	
	RU-PNPI_DATADISK	0 %	delta	
	RU-PROTVINO-IHEP_DATADISK	1.2 %	bravo	
	TECHNION-HEP_DATADISK	0 %	delta	
	TR-16-ILAKHIM_DATADISK	0.6 %	charlie	
UK	WEIZMANN-LCG2_DATADISK	1.2 %	bravo	registered: 13 ready: 13
	UKI-LT2-QMUL_DATADISK	2.4 %	alpha	
	UKI-LT2-RHUL_DATADISK	1.2 %	bravo	
	UKI-LT2-UCL-HEP_DATADISK	0 %	delta	
	UKI-NORTHGRID-LANGS-HEP_DATADISK	2.4 %	alpha	
	UKI-NORTHGRID-LIV-HEP_DATADISK	1.2 %	bravo	
	UKI-NORTHGRID-MAN-HEP_DATADISK	2.4 %	alpha	
	UKI-NORTHGRID-SHEP-HEP_DATADISK	1.2 %	bravo	
	UKI-SCOTGRID-ECDF_DATADISK	2.4 %	alpha	
	UKI-SCOTGRID-GLASGOW_DATADISK	2.4 %	alpha	
	UKI-SOUTHGRID-BHAM-HEP_DATADISK	1.2 %	bravo	
	UKI-SOUTHGRID-CAM-HEP_DATADISK	1.2 %	bravo	
	UKI-SOUTHGRID-OX-HEP_DATADISK	2.4 %	alpha	
	UKI-SOUTHGRID-RALPP_DATADISK	1.2 %	bravo	
US	AGLT2_DATADISK	0.6 %	charlie	registered: 7 ready: 7
	ILLINOISHEP_DATADISK	0 %	delta	
	MWT2_DATADISK	2.4 %	alpha	
	NET2_DATADISK	2.4 %	alpha	
	OU-OCHEP-SWT2_DATADISK	0 %	delta	
	SLACKRDX_DATADISK	2.4 %	alpha	
	SWT2-CPB_DATADISK	2.4 %	alpha	



Distribuzione dei dati

Evoluzione del Computing Model

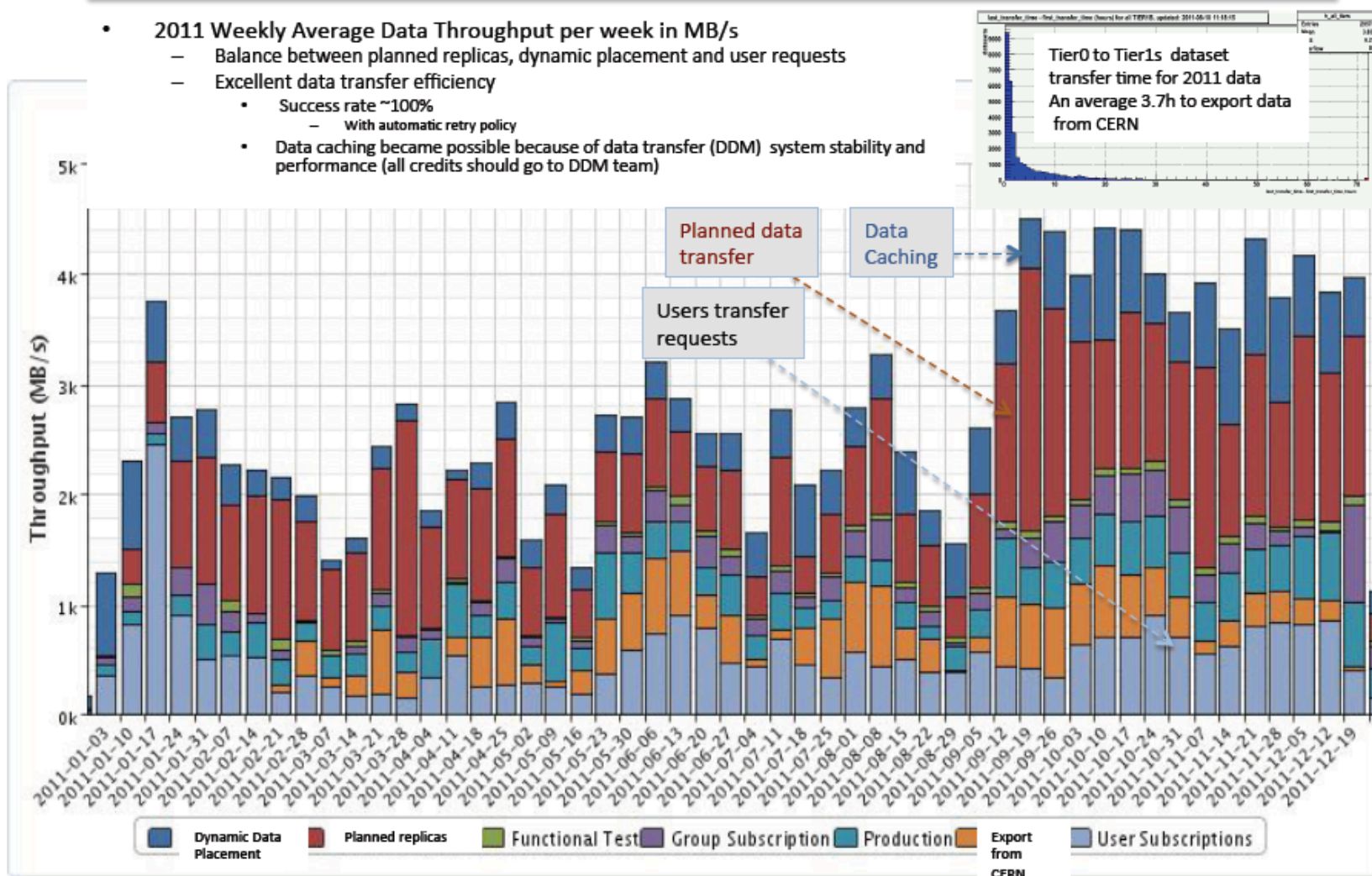


Distribuzione dei dati



- Statica (planned) = distribuzione predefinita secondo share fissati
- Dinamica (data caching) = distribuzione in base alla popolarità dei dati

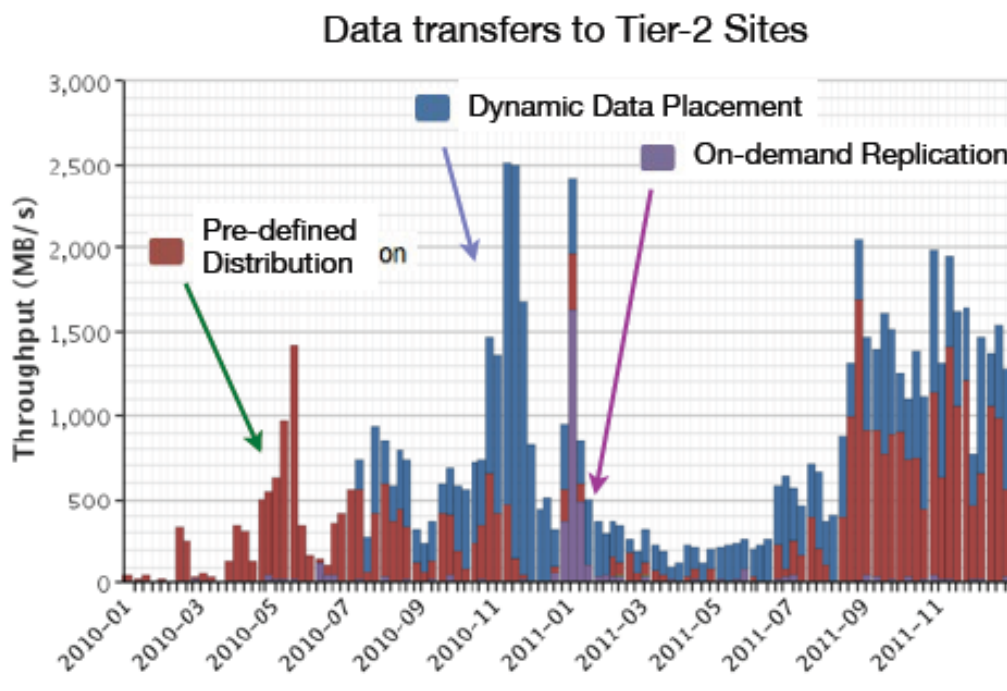
- 2011 Weekly Average Data Throughput per week in MB/s
 - Balance between planned replicas, dynamic placement and user requests
 - Excellent data transfer efficiency
 - Success rate ~100%
 - With automatic retry policy
 - Data caching became possible because of data transfer (DDM) system stability and performance (all credits should go to DDM team)



Distribuzione dei dati



- Nel 2011 si è cercato il giusto rapporto tra il volume di dati trasferiti dinamicamente e staticamente per permettere che una frazione significativa dell'analisi fosse svolta nei Tier2 insieme ad un uso saggio del disco
- Le comunità di utenti fanno capo soprattutto ai Tier2 (cpu e disco dedicati, accesso diretto), era dannoso concentrare l'analisi solo ai Tier1 (inizi 2011)



Replica dei dati prestabilita:

- Tier1, replica per ridondanza (consolidamento), copia primaria
- Tier2, replica per l'analisi, primaria
- Tier2, replica extra per l'analisi, copia secondaria

Determinata dallo share della cloud (Tier1) e dalla classificazione (Tier2)

Replica dinamica dei dati

- Tier1 e Tier2, basata sulla popolarità dei dati, copia secondaria

Determinata dall'utilizzo dei siti

Dati disponibili in Italia



- DATADISK
 - Area di storage centrale (pledged)
 - contiene sia i dati preplaced prodotti centralmente che quelli prodotti dal PD2P, circa il 10% del totale. I secondari vengono cancellati quando non utilizzati
- GROUPDISK
 - Aree gestite dai gruppi (pledged), dati sempre più utilizzati per l'analisi
 - CNAF: SUSY e TOP
 - Milano: MS e EGAMMA
 - Napoli: HIGGS e TRIGGER
 - Roma1: HIGGS e MUONI
- LOCALGROUPDISK
 - Area di storage (quasi) permanente per gli utenti italiani
 - Uso da massimizzare sfruttando l'accesso diretto dai *Tier3*, root o Proof on Demand, o con job di griglia (prun)

- Sottoscrizione automatica dell'output dell'analisi sul proprio LOCALGROUP quando si esegue l'analisi in griglia, in qualsiasi sito (solo per i dati da conservare)
 - option "--DESTSE=INFN-NAPOLI-ATLAS_LOCALGROUPDISK "
- Da pochissimo disponibile un tool, basato su dq2, che manda una mail quando un dataset arriva nel sito specificato (grazie a Manoj)
 - utilissimo per minimizzare la latenza dell'analisi



Accesso ai dati



- Accesso in GRID
 - Pathena, Prun, Proof on Demand
 - Analisi diretta o skimming su AOD, ESD, D3PD, NTUP
 - Mette a disposizione un numero molto maggiore di risorse sul singolo sito o, parallelizzando l'attività, su più centri di calcolo distribuiti a livello geografico
 - Necessario per tutte le attività su grandi sample di dati o urgenti
- Accesso diretto
 - Root, Proof, athena per sviluppo e test di codice
 - Analisi finale
 - Accesso diretto dalle CPU *à la Tier3* ai dati nello storage del Tier2. Nativo nei siti Storm (MI) o, grazie al protocollo XROOTD adottato alla fine dell'anno scorso, nei siti DMP (LNF, NA, RM1)
 - Evita la duplicazione delle NTUP per l'analisi finale localmente nel Tier3

Tipica analisi in 2 step:

1. Skimming in GRID su dati di gruppo con produzione di NTUP light salvate in LOCALGROUP
2. Analisi interattiva con ROOT/PROOF o in GRIGLIA con PoD sui dati
vedi esempio analisi $H \rightarrow ZZ$ di Valerio

Proof on Demand (PoD)



Elisabetta Vilucchi e
Roberto De Nardo (LNF)

- **Proof** – tool che parallelizza l'analisi con ROOT su diversi core della stessa macchina (Proof Light) o su un cluster di nodi
 - **Demand** – possibilità di usare, a richiesta, i nodi di una farm destinata prevalentemente ad altri scopi (per esempio Tier2/3 in GRID)
- L'utente lancia il master PoD sulla propria macchina che apre un canale verso il cluster e consente all'utente di richiedere l'allocazione di un certo numero di job slot
 - A questo punto l'utente invia la richiesta che si traduce in un numero di job in coda pari al numero di job slot richieste
 - L'utente può controllare l'allocazione delle risorse che gli appariranno disponibili non appena i job vanno in stato di running. A questo punto l'utente può sottomettere l'analisi con Proof sui nodi di calcolo in cui compare un job PoD running
 - Proof, in base ai nodi a disposizione, suddivide l'analisi in un certo numero di job paralleli

Proof on Demand (PoD)



Elisabetta Vilucchi e
Roberto De Nardo (LNF)

Proof on Demand su un cluster Grid: Tier2 o Tier3

- E' stato sviluppato un plugin di PoD per gLite, che da la possibilità agli utenti di attivare un cluster Proof “on demand” su una farm in Grid con middleware gLite
- Gli utenti, connettendosi ad una UI, possono lanciare PoD e riservare un certo numero di nodi sulla farm di un Tier2/3
- La gestione delle risorse e' simile a quella del cluster locale e il codice per il setup di PoD viene fatto direttamente da cvmfs, disponibile ormai nella maggior parte dei siti di ATLAS
- Da ultimare, prevista per un prossimo sviluppo, un plugin di PoD per Panda, il sistema ufficialmente riconosciuto dalla collaborazione per fare analisi, in modo tale che sia possibile visualizzare nell'accounting dell'esperimento anche i job di analisi fatti con questo plugin

Proof on Demand (PoD)



Elisabetta Vilucchi e
Roberto De Nardo (LNF)

- Da domani (18-19/04) si terrà un tutorial ai Laboratori Nazionali di Frascati per illustrare l'utilizzo di Proof, PoD e del nuovo plugin per gLite:
<http://agenda.infn.it/conferenceDisplay.py?confId=4933>
- Per partecipare è sufficiente possedere un certificato digitale, essere iscritti alla VO ATLAS ed avere un account al Cern
- Molto utile per tutti coloro che fanno analisi

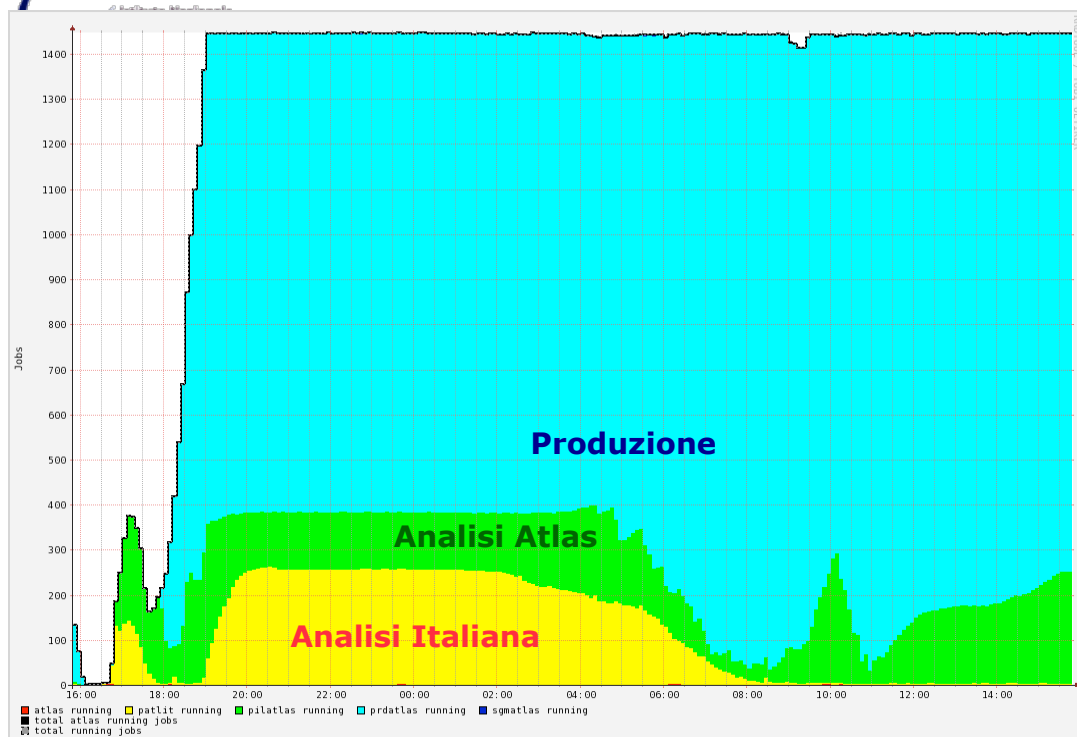


Utilizzo risorse Italiane



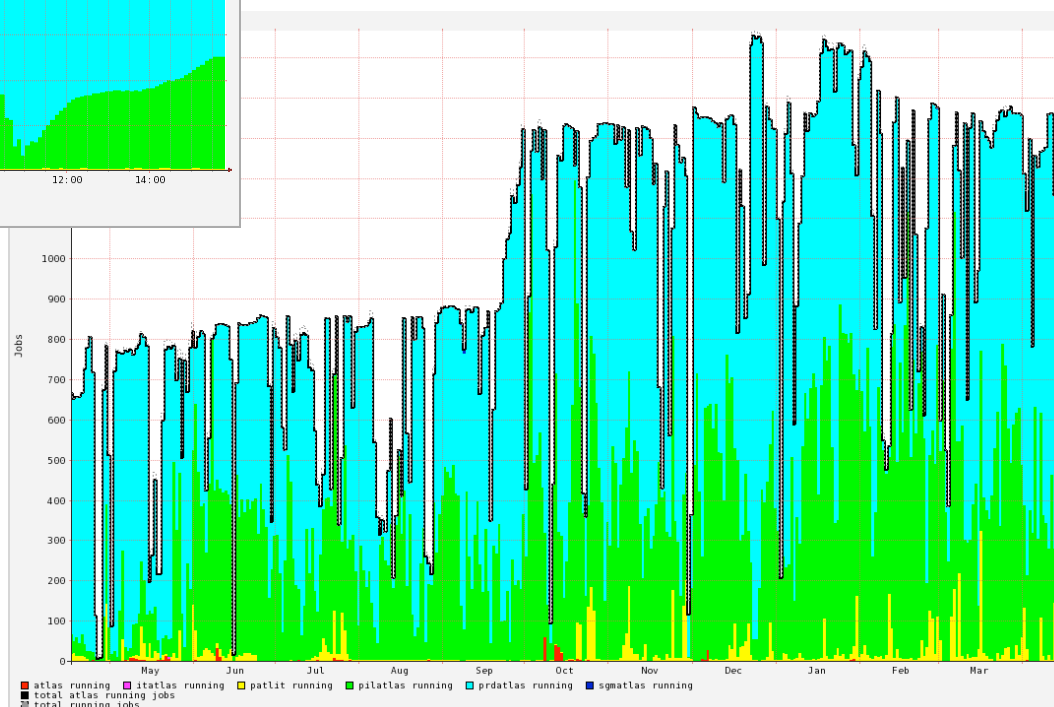
Utilizzo efficiente delle risorse: priorità degli utenti italiani nei siti Grid di ATLAS Italia

- **Gli utenti italiani possono utilizzare uno share dedicato dei siti**
 - Macchine fisicamente riservate o slot ad alta priorità
 - Gli utenti italiani, identificati automaticamente da Panda tramite la presenza del gruppo VOMS **/atlas/it** nel loro proxy accedono agli share dedicati, senza alcuna operazione da parte loro
 - I job sottomessi ad un sito italiano da un utente italiano che ha il gruppo **/atlas/it** partiranno prima degli altri job di analisi, anche se accodati in momenti successivi
 - Per assicurarsi di avere l'accesso allo share di **/atlas/it** è sufficiente controllare nel momento della registrazione in VOMS di richiedere il gruppo **/atlas/it**, **anche se il certificato non proviene dall'INFN, a patto che si faccia parte della comunità italiana**
 - Nel caso in cui non si possenga il gruppo si può farne richiesta anche dopo essersi registrati



Job running in un Tier2:

← Ultima settimana
↓ Ultimo anno



- L'analisi viene svolta efficacemente nei Tier2 Italiani
- Le risorse dedicate (dedicabili) sono significative
- Permettono ai job italiani di andare in run più velocemente senza essere accodati agli altri
- **Non solo analisi, anche prod MC**

Produzione Monte Carlo



ATLAS accetta le produzioni Monte Carlo private (G4, digi e reco) su risorse *unpledged*

- **Additional Production**

- Produzione di gruppo in aggiunta a quelle già concordate
- Deve essere approvata dal Physics Coordinator specificando nella richiesta che si tratta di produzione Addizionale
 - Sanity check contro richieste poco motivate
- Eseguita centralmente (Prod sys) ma i dati non sono distribuiti dal DDM
 - Solo il PD2P può spostare questi dati
 - Il team di simulazione invia al sito specificato che offre le risorse

- **Private Production**

- User production diverse da quelle ufficiali di ATLAS. Non vista di buon occhio
- Eseguita individualmente
- Richiesta la validazione confrontandola con un sample di riferimento eseguito centralmente

Produzione Monte Carlo



Qualche numero

- *Simulation Time = 3300 HepSpec06 sec*
- *Reconstruction Time = 830 HepSpec06 sec*
 - valori a 8 TeV ottenuti dalla media dei tempi a 7 TeV sull'intero set di canali (e # totale di eventi) scalato a 8 TeV

Potenza di calcolo a disposizione:

- Valor medio di un singolo core = 10 HepSpec06
 - 1 evento richiede ~ 7 minuti per sim e reco
 - Potenza di calcolo totale nei Tier2 (estate 2012) = ~52 kHS06
 - Stima potenza di calcolo per attività italiane = ~10 kHS06
 - Dopo aver sottratto: WLCG Pledges = 26.6 kHS06 e varie = ~2.5 kHS06 e considerato: efficienza = 85%, risorse non disponibili = 5%
 - da usare per analisi e produzione MC con percentuali da decidere
- 10⁵ eventi al giorno e 3x10⁶ al mese (nell'ipotesi share 50%)**

Produzione Monte Carlo



10^5 eventi al giorno e 3×10^6 al mese (nell'ipotesi share 50%)

Attenzione, è una novità, quindi è da testare e verificare

- Verificare se i gruppi di fisica sono d'accordo, soprattutto nel caso di produzioni private, e le procedure da seguire in tal caso
- Dal punto di vista dei siti, richiede la pubblicazione delle CPU totali e pledged in Panda e il test che la procedura di assegnazione dei task e delle risorse dedicate funzioni correttamente
- Siamo pronti a collaborare con i gruppi italiani e i responsabili dell'analisi per definire la lista di produzioni MC e il corretto peso rispetto all'analisi e collaborare con ATLAS per i test