



TempoQuestTM

Acceleration of WRF on the GPU

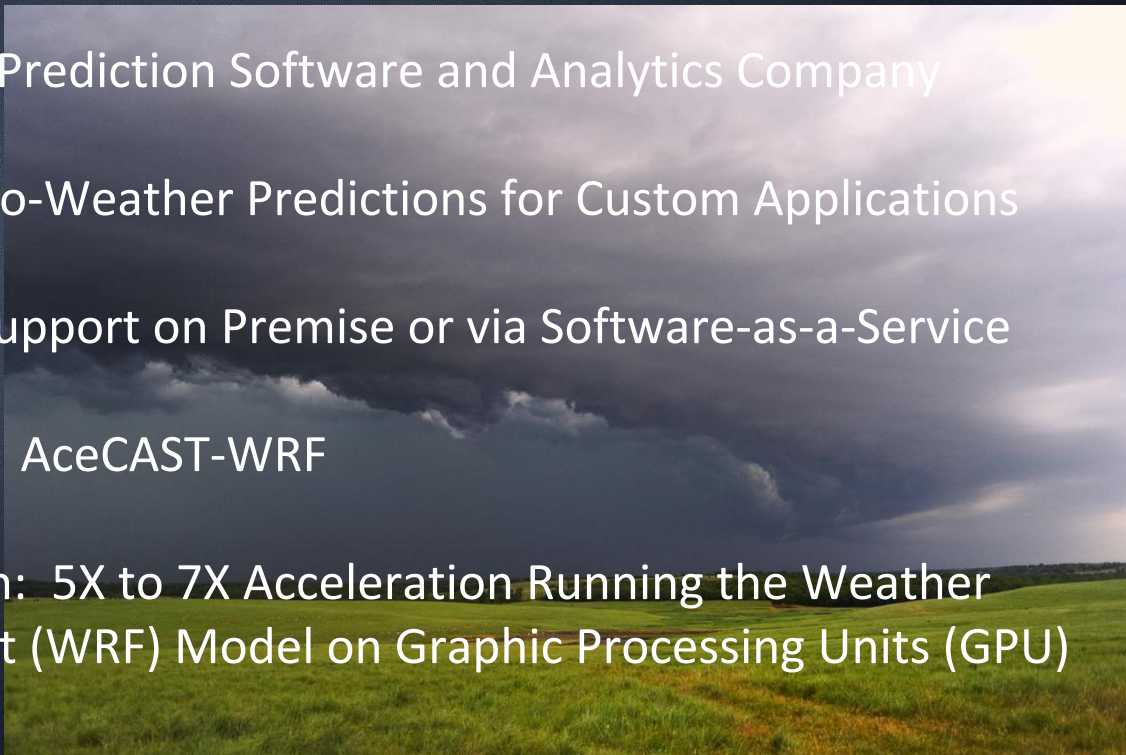
**Daniel Abdi, Sam Elliott, Iman Gohari
Don Berchoff, Gene Pache, John Manobianco**

TempoQuest
1434 Spruce Street
Boulder, CO 80302
720 726 9032
TempoQuest.com

THE WORLD'S FASTEST MOST PRECISE FORECASTS

Our Product

- TQI is a Weather Prediction Software and Analytics Company
- We Produce Micro-Weather Predictions for Custom Applications
- We Deliver and Support on Premise or via Software-as-a-Service
- Flagship Product: AceCAST-WRF
- The Breakthrough: 5X to 7X Acceleration Running the Weather Research Forecast (WRF) Model on Graphic Processing Units (GPU)



Our approach to Re-factoring

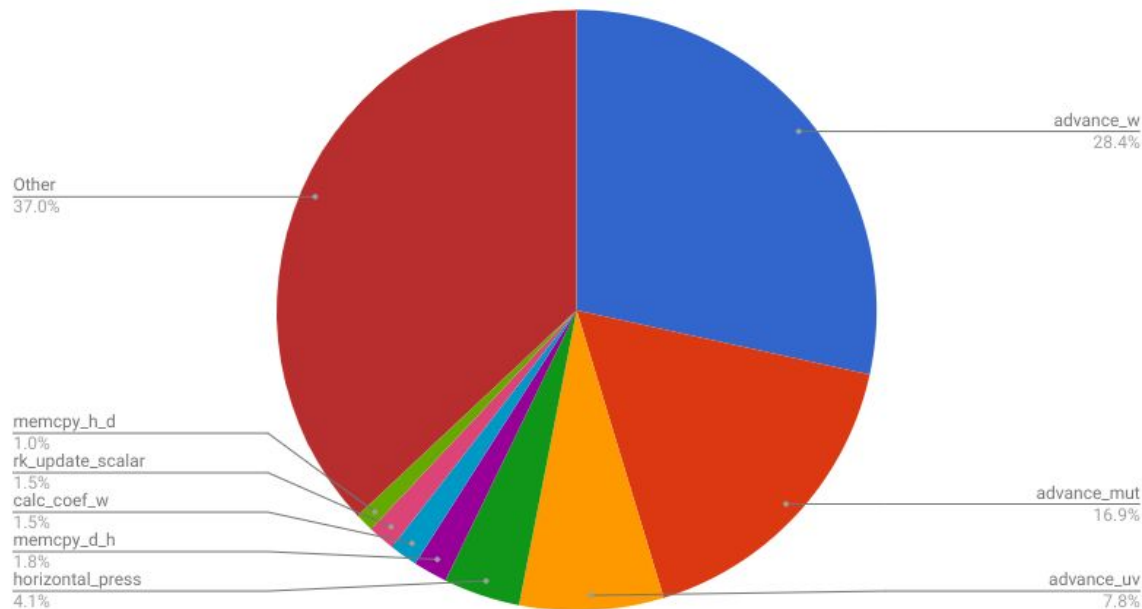
- WRF ported to run entirely on the GPU
- Profile and optimize most time consuming parts
- Avoid/minimize data transfer to/from GPU
- Leverage WRF registry to produce GPU code
- Pack halo data on GPU and send via infiniband
- Process multiple tiles and columns in a kernel

Our approach to Re-Factoring

- Two branches : hybrid CPU + GPU vs pure GPU
- 7x difference in speedup between those two
- “Premature optimization is the root of all evil”
- Parallelize->Profile->Optimize->Rewrite & Repeat

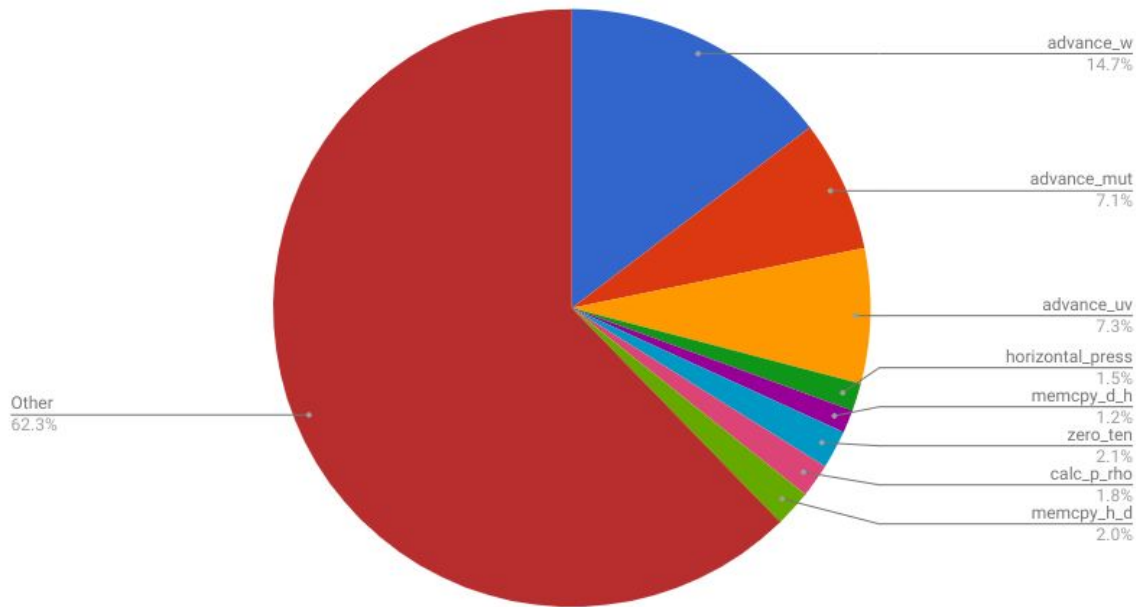
Profile on P100 GPU - Before Optimization

Wrf dynamics profile: Before Optimization



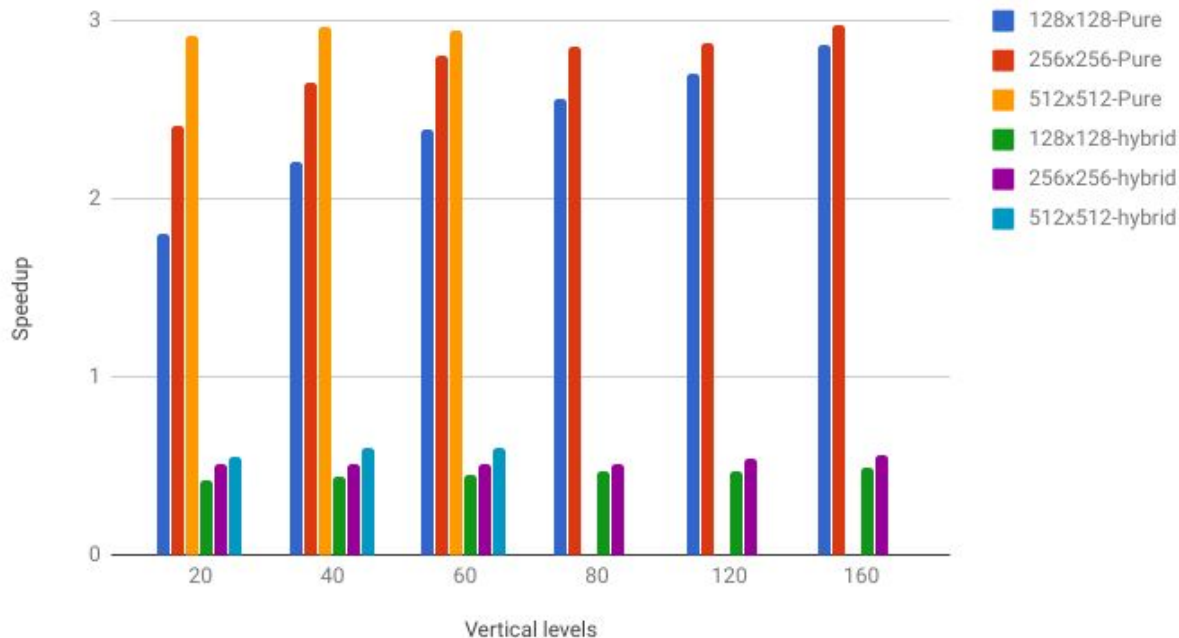
Profile on P100 GPU - After Optimization

Wrf dynamics profile: After Optimization



Cost of data transfer- P100 GPU + Haswell CPU

GPU speedup vs 1-core CPU on Pure GPU and Hybrid CPU-GPU modes



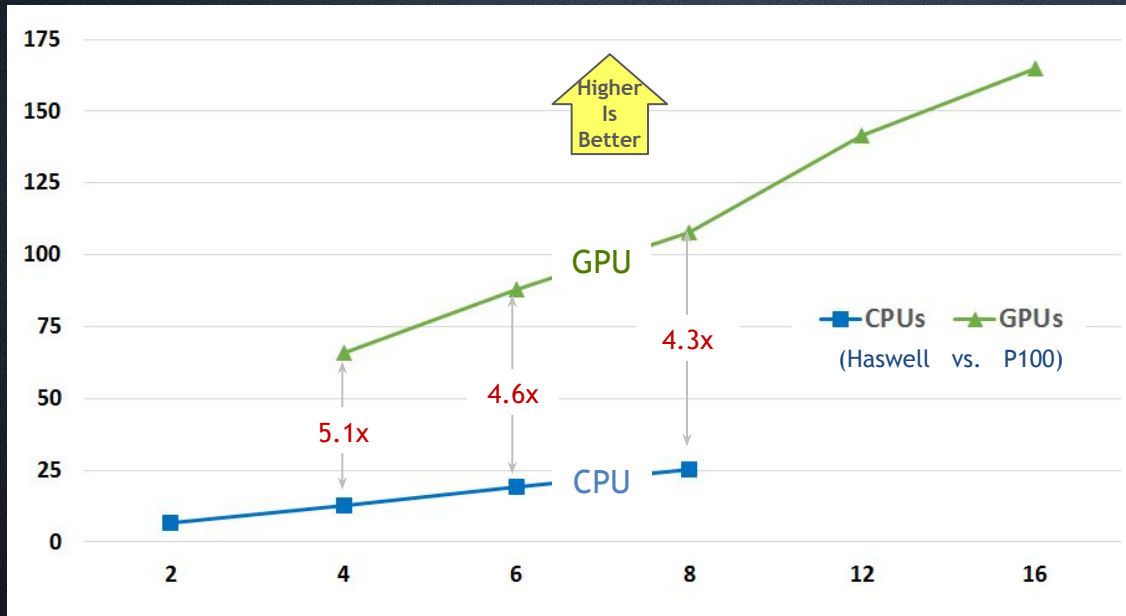
Results: GPU WRF Strong Scaling for CONUS 2.5 km

~5x Speedup Full Model: 4 x P100 vs. 4 x HSW

(1 x node)

(2 x nodes)

Performance [MM grid points / sec]



Number of Processors (CPUs or GPUs)

CONUS 2.5 km Case on PSG Cluster - 4 nodes

Source: TQI – Abdi; Apr 18

- Based on WRF 3.8.1 trunk
- 1501 x 1201 grid x 35 levels
- Total 60 time steps, SP run
- **Physics option modified:**
 - WSM6
 - Radiation *off*
 - 5-layer TDS
- All WRF runs single precision
- PSG cluster node configuration:
 - 2 CPUs, 16 cores each
 - 4 x P100 GPUs
 - Or 4 x V100 GPUs
- CPU-only 1 MPI task each core
- CPU+GPU 1 MPI task per GPU

CONUS 2.5km Source: http://www2.mmm.ucar.edu/wrf/bench/benchdata_v3911.html (Note "Physics options modified" in side bar)

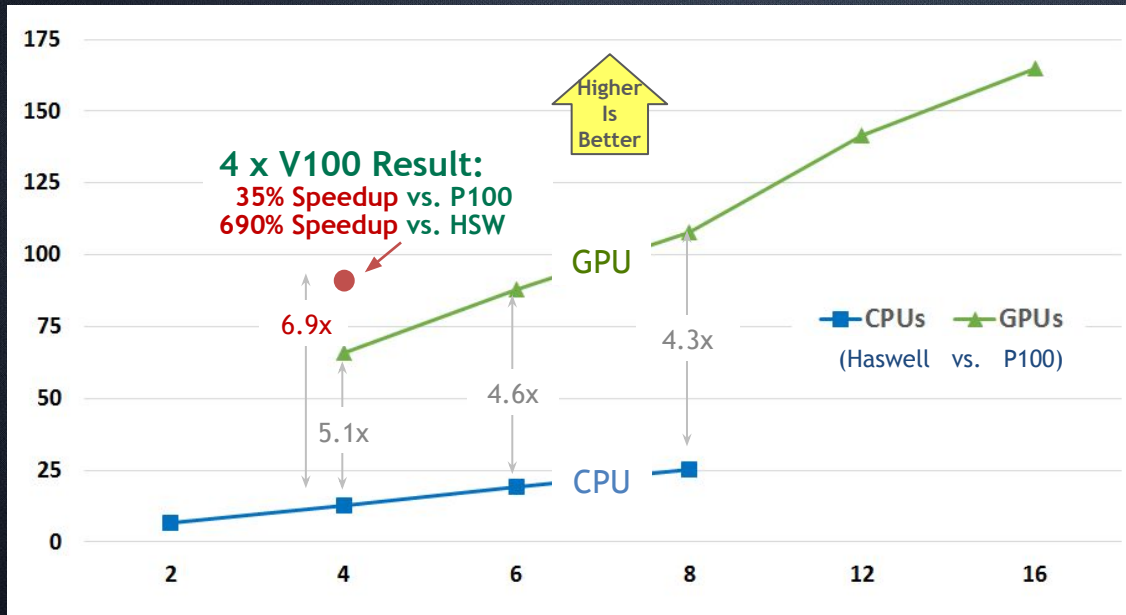
Results: GPU WRF Strong Scaling for CONUS 2.5 km

~7x Speedup Full Model: 4 x V100 vs. 4 x HSW

(1 x node)

(2 x nodes)

Performance [MM grid points / sec]



Number of Processors (CPUs or GPUs)

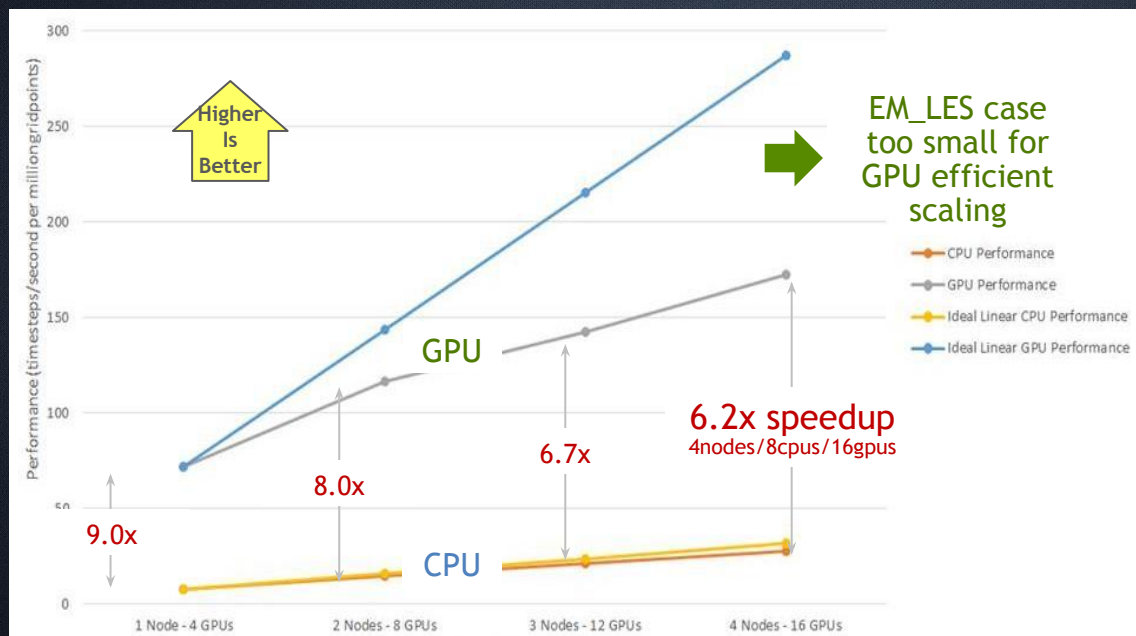
CONUS 2.5 km Case on PSG Cluster - 4 nodes

Source: TQI – Abdi; Apr 18

- Based on WRF 3.8.1 trunk
- 1501 x 1201 grid x 35 levels
- Total 60 time steps, SP run
- **Physics option modified:**
 - WSM6
 - Radiation *off*
 - 5-layer TDS
- All WRF runs single precision
- PSG cluster node configuration:
 - 2 CPUs, 16 cores each
 - 4 x P100 GPUs
 - Or 4 x V100 GPUs
- CPU-only 1 MPI task each core
- CPU+GPU 1 MPI task per GPU

Results: GPU WRF Strong Scaling for EM_LES

~5x Speedup: 4 x P100 vs. 4 x HSW



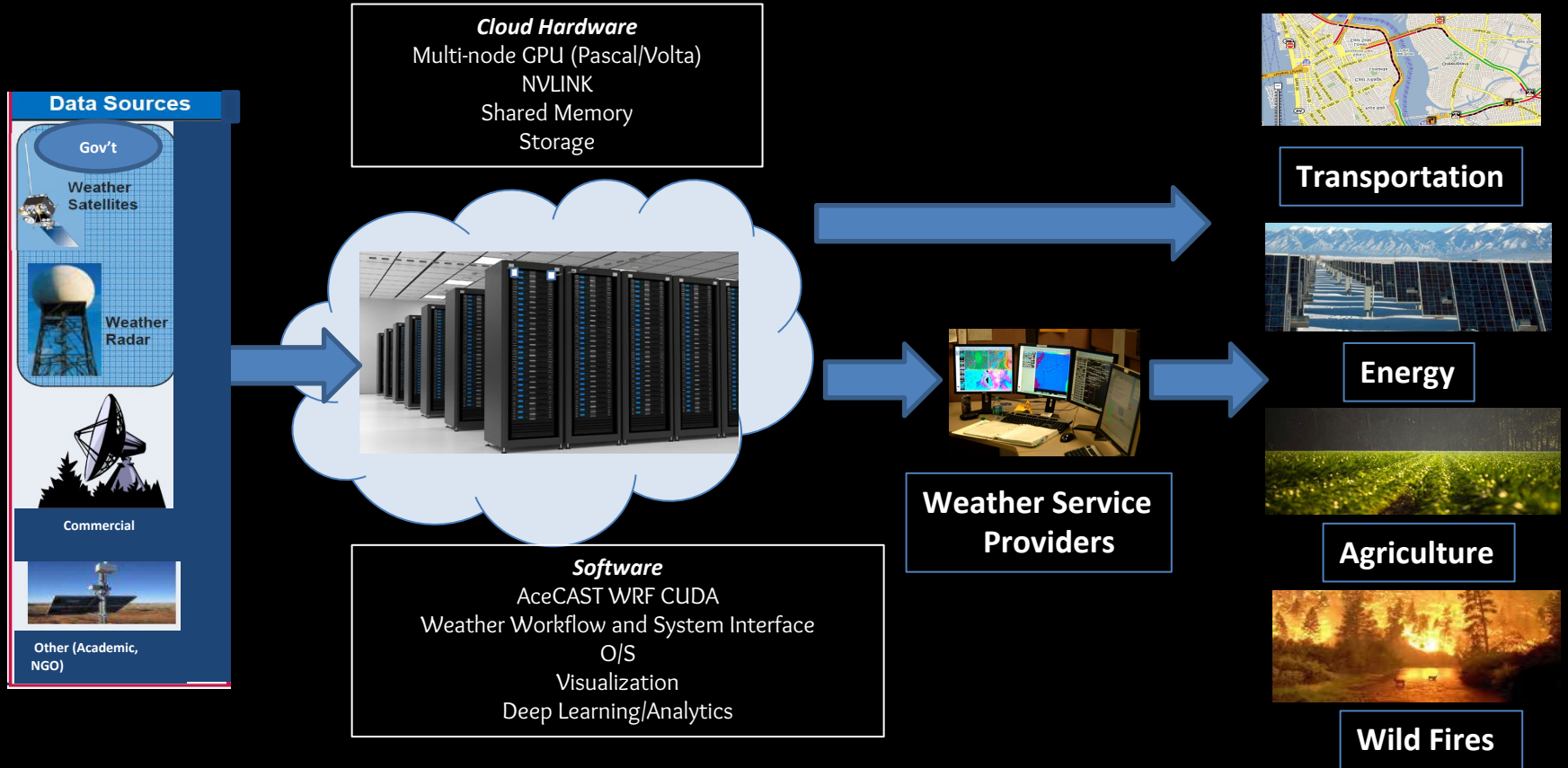
Number of Nodes

Results for EM_LES Case on PSG - 4 nodes

Source: TQI – Abdi; Dec 18

- Based on WRF 3.8.1 trunk
- 1024 x 1024 grid x 60 levels
- Physics options:
 - Kessler
 - Mostly dycore time
- PSG cluster nodes:
 - 2 CPUs, 16 cores each
 - 4 x P100 GPUs
- CPU-only MPI task each core
- CPU+GPU MPI task per GPU

TempoQuest Systems Architecture



Conclusions

- TQI is a micro-weather prediction company with the goal of accelerating WRF by up to 10x using NVIDIA GPUs
- We had a breakthrough with acceleration of end-to-end WRF runs by 5x to 7x
- We deliver on-premise or software-as-service on the cloud
- Future goal: we feel the need for more speed ...