

GPU Developments in Atmospheric Sciences

Stan Posey, HPC Program Manager, ESM Domain, NVIDIA (HQ), Santa Clara, CA, USA
David Hall, Ph.D., Sr. Solutions Architect, NVIDIA, Boulder, CO, USA



TOPICS OF DISCUSSION

- **NVIDIA HPC UPDATE**
- **GPU RESULTS FOR NUMERICAL MODELS: IFS, WRF**
- **GPU BENEFITS FOR AI MODELS**

NVIDIA Company Update

HQ in Santa Clara, CA, USA

FY18 Rev ~\$10B USD

GTC 2019 Mar 18-22

www.gputechconf.com/present/call-for-submissions

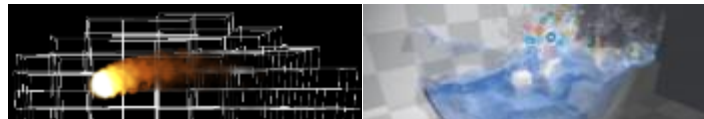


New NVIDIA HQ “Endeavor”

NVIDIA Core Markets



GPU (HPC) Computing



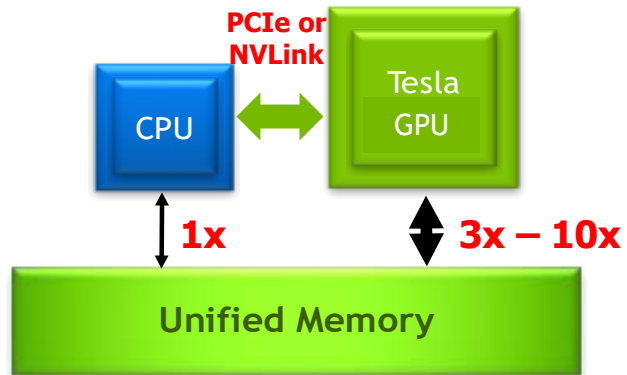
Computer Graphics



Artificial Intelligence

NVIDIA GPU Introduction

GPU Introduction



- Co-processor to the CPU
- Threaded Parallel (SIMT)
- CPUs: x86 | Power
- HPC Motivation:
 - Performance
 - Efficiency
 - Cost Savings

#1 ORNL Summit: 4,600 nodes; **27,600 GPUs**



NVIDIA GPUs in the Top 10 of Current Top500

"It's somewhat ironic that training for deep learning probably has more similarity to the HPL benchmark than many of the simulations that are run today..."

- - Kathy Yelick, LBNL



- Top500 #1: ORNL Summit
 - RMAX (HPL) = **122 PF**
 - IBM Power + GPU system
- Top #1's in 3 regions with GPUs:
 - USA - ORNL Summit
 - Europe - CSCS Piz Daint
 - Japan - AIST ABCI
- Total 5 of Top7 with GPUs
- #1 Summit vs. #7 Titan (2012)
 - ~7x more performance
 - ~Same power consumption

Rank	System	Cores	Rmax (TFlop/s)	Rpeak (TFlop/s)	Power (kW)
1	Summit - IBM Power System AC922, IBM POWER9 22C 3.07GHz, <u>NVIDIA Volta GV100</u> , Dual-rail Mellanox EDR Infiniband , IBM DOE/SC/Oak Ridge National Laboratory United States	2,282,544	122,300.0	187,659.3	8,806
2	Sunway TaihuLight - Sunway MPP, Sunway SW26010 260C 1.45GHz, Sunway , NRCPC National Supercomputing Center in Wuxi China	10,649,600	93,014.6	125,435.9	15,371
3	Sierra - IBM Power System S922LC, IBM POWER9 22C 3.1GHz, <u>NVIDIA Volta GV100</u> , Dual-rail Mellanox EDR Infiniband , IBM DOE/NNSA/LLNL United States	1,572,480	71,610.0	119,193.6	
4	Tianhe-2A - TH-IVB-FEP Cluster, Intel Xeon E5-2692v2 12C 2.2GHz, TH Express-2, Matrix-2000 , NUDT National Super Computer Center in Guangzhou China	4,981,760	61,444.5	100,678.7	18,482
5	AI Bridging Cloud Infrastructure (ABCI) - PRIMERGY CX2550 M4, Xeon Gold 6148 20C 2.4GHz, <u>NVIDIA Tesla V100 SXM2</u> , Infiniband EDR , Fujitsu National Institute of Advanced Industrial Science and Technology [AIST] Japan	391,680	19,880.0	32,576.6	1,649
6	Piz Daint - Cray XC50, Xeon E5-2690v3 12C 2.6GHz, Aries interconnect , <u>NVIDIA Tesla P100</u> , Cray Inc. Swiss National Supercomputing Centre [CSCS] Switzerland	361,760	19,590.0	25,326.3	2,272
7	Titan - Cray XK7, Opteron 6274 16C 2.200GHz, Cray Gemini interconnect, <u>NVIDIA K20x</u> , Cray Inc. DOE/SC/Oak Ridge National Laboratory United States	560,640	17,590.0	27,112.5	8,209

Feature Progression in NVIDIA GPU Architectures

	V100 (2017)	P100 (2016)	K40 (2014)
Double Precision TFlop/s	7.5 1.4x - <u>5.4x</u>	5.3 3.8x	1.4
Single Precision TFlop/s	15.0 1.4x - <u>3.5x</u>	10.6 2.5x	4.3
Half Precision TFlop/s	120 (DL) ~6x	21.2	n/a
Memory Bandwidth (GB/s)	900 1.25x - <u>3.1x</u>	720 2.5x	288
Memory Size	16 or 32GB 2.00x	16GB 1.33x	12GB
Interconnect	NVLink: Up to 300 GB/s PCIe: 32 GB/s	NVLink: 160 GB/s PCIe: 32 GB/s	PCIe: 16 GB/s
Power	250W - 300W	300W	235W

V100 Availability



NVIDIA GPU Server Node System for Data Centers

Tesla HGX-2

Building Block for Large-scale Systems

16 Tesla V100 GPUs | 512 GB Memory | 14.4 TB/s BW

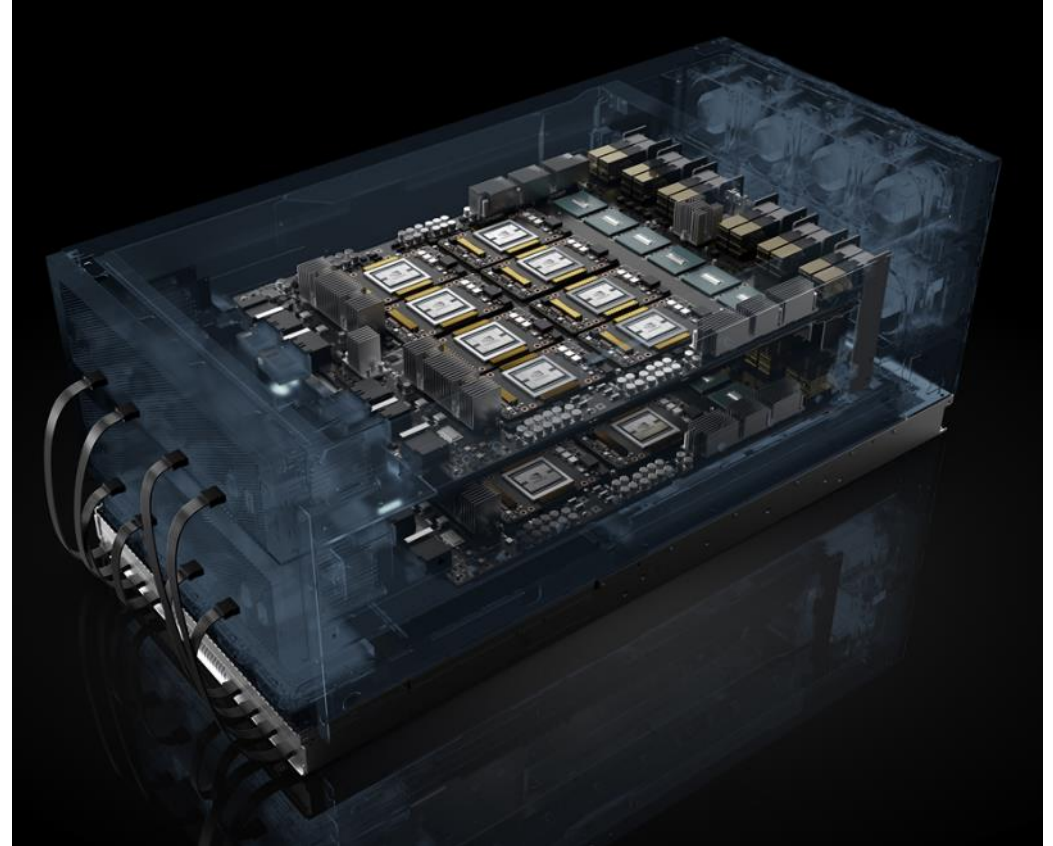
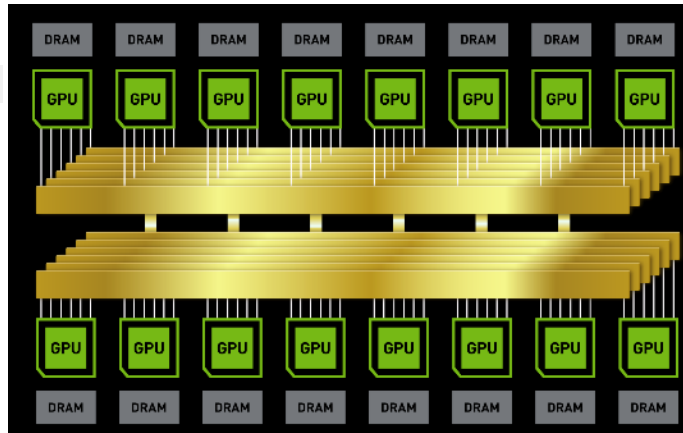
Multi-precision Computing

2 PFLOPS AI | 250 TFLOPS FP32 | 125 TFLOPS FP64

16 NVIDIA Tesla V100 GPUs

12 NVIDIA NVSwitches

Direct GPU-to-GPU Connection
Between All 16 GPUs



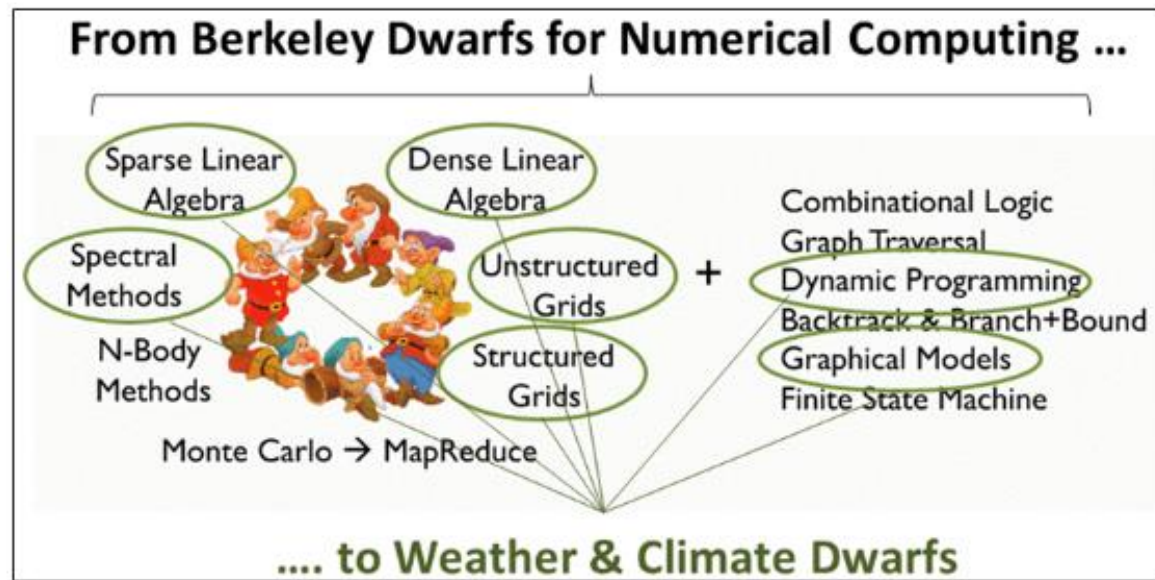
GPU Collaborations for Atmospheric Models

	Model	Organizations	Funding Source	
	E3SM-Atm, SAM	DOE: ORNL, SNL	E3SM, ECP	 
	MPAS-A	NCAR, UWyo, KISTI, IBM	WACA II	 
	FV3/UFS	NOAA	SENA	
	NUMA/NEPTUNE	NPS, US Naval Res Lab	NPS	
	IFS	ECMWF	ESCAPE	
	GungHo/LFRic	MetOffice, STFC	PSyclone	
	ICON	DWD, MPI-M, CSCS, MCH	PASC ENIAC	
	KIM	KIAPS	KMA	
	COSMO	MCH, CSCS, DWD	PASC GridTools	
	WRFg	NVIDIA, NCAR	NVIDIA	
	AceCAST-WRF	TempoQuest	Venture backed	

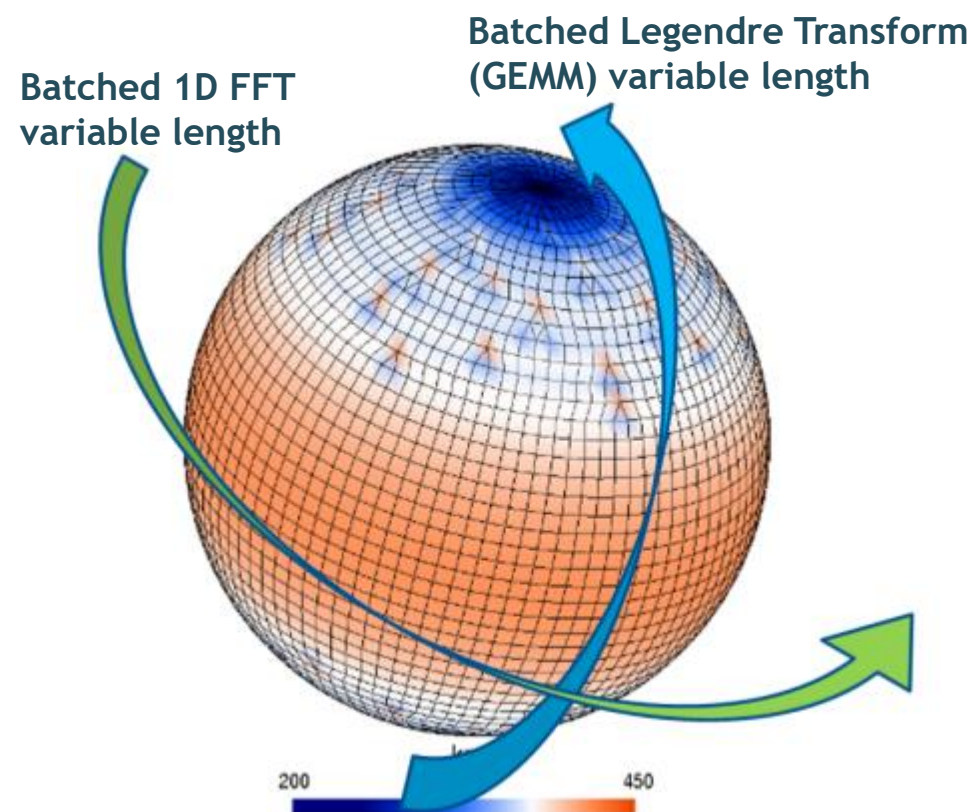
TOPICS OF DISCUSSION

- NVIDIA HPC UPDATE
- GPU RESULTS FOR NUMERICAL MODELS: **IFS, WRF**
- GPU BENEFITS FOR AI MODELS

ESCAPE Development of Weather & Climate Dwarfs



NVIDIA-Developed Dwarf: Spectral Transform - Spherical Harmonics (SH)

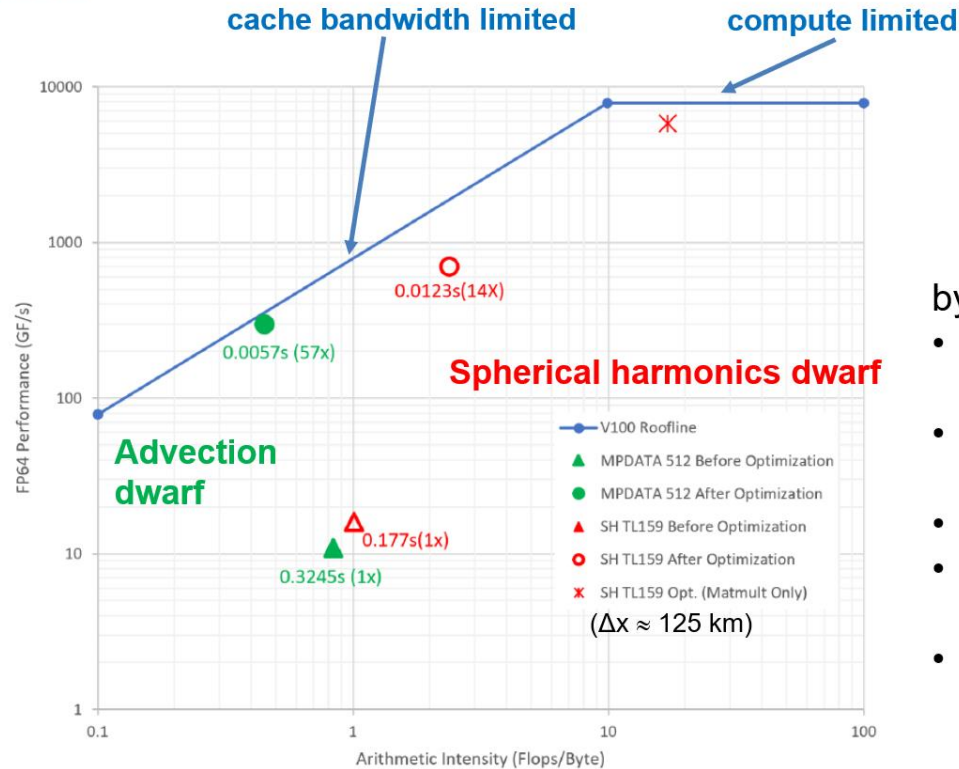


GPU-based Spectral Transform Approach:

- For all vertical Layers
 - 1D FFT along all latitudes
 - GEMM along all longitudes
- NVIDIA-based libraries for FFT and GEMM
- OpenACC directives

GPU Performance of SH Dwarf in IFS - Single-GPU

Hybrid Computing – single GPU



by:

- exposing parallelism in loops for OpenACC mapping
- Kernel optimization by memory mapping
- exploiting CUDA BLAS features
- minimizing data allocation and movement
- (calling C/CUDA from PGI Fortran)

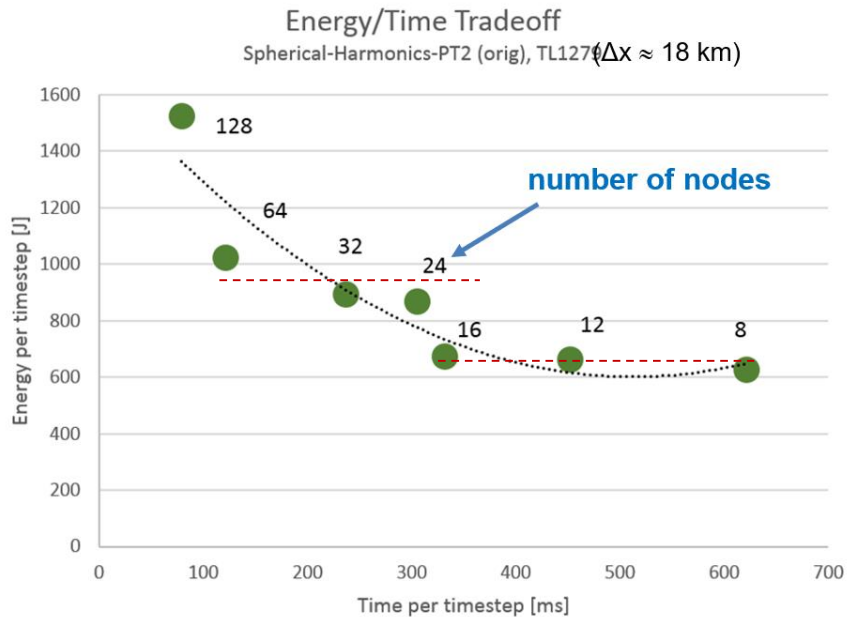
From “ECMWF Scalability Programme”

Dr. Peter Bauer,
UM User Workshop,
MetOffice, Exeter, UK
15 June 2018

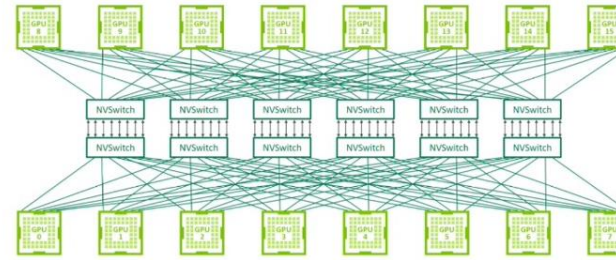
- Single V100 GPU improved SH dwarf by 14x
- Single V100 GPU improved MPDATA dwarf by 57x

GPU Performance of SH Dwarf in IFS - Multi-GPU

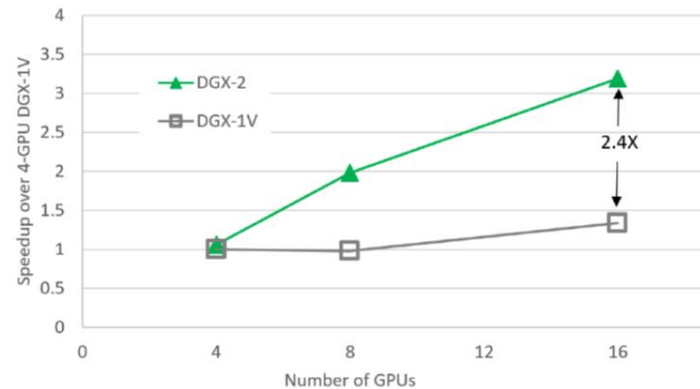
Hybrid Computing – multiple GPU



NVIDIA NGX-2 with NVSwitch



Spherical Harmonics Dwarf TCO639 Test Case ($\Delta x \approx 18$ km)
DGX-2 vs DGX-1V



From “ECMWF Scalability Programme”

Dr. Peter Bauer,
UM User Workshop,
MetOffice, Exeter, UK
15 June 2018

- Results of Spherical Harmonics Dwarf on NVIDIA DGX System
- Additional 2.4x gain from DGX-2 NVSwitch for 16 GPU systems

WRF GPU Status - WRFg - Sep 2018



WRFg



- Based on ARW release 3.7.1
- Limited physics options:
 - Enough for full model on GPU
 - Further development ongoing
- Community release:
 - Freely available objects/exe
 - Restricted source availability
- Availability during Q1 2019:
 - Benchmark tests – NOW
- Support and roadmap – TBD

WRFg Physics Options (10)

Microphysics

Option

WSM5	4
WSM6	6
Thompson	8

Radiation

RRTM (lw), Dudhia (sw)	1
RRTMG_fast (lw,sw)	24

Planetary boundary layer

YSU	1
-----	---

Surface layer

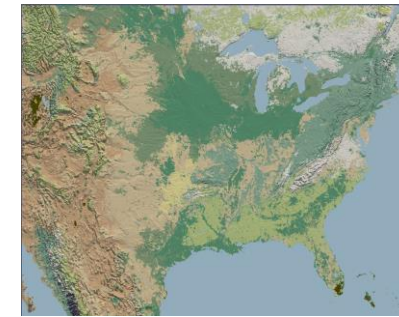
Revised MM5	1
-------------	---

Land surface

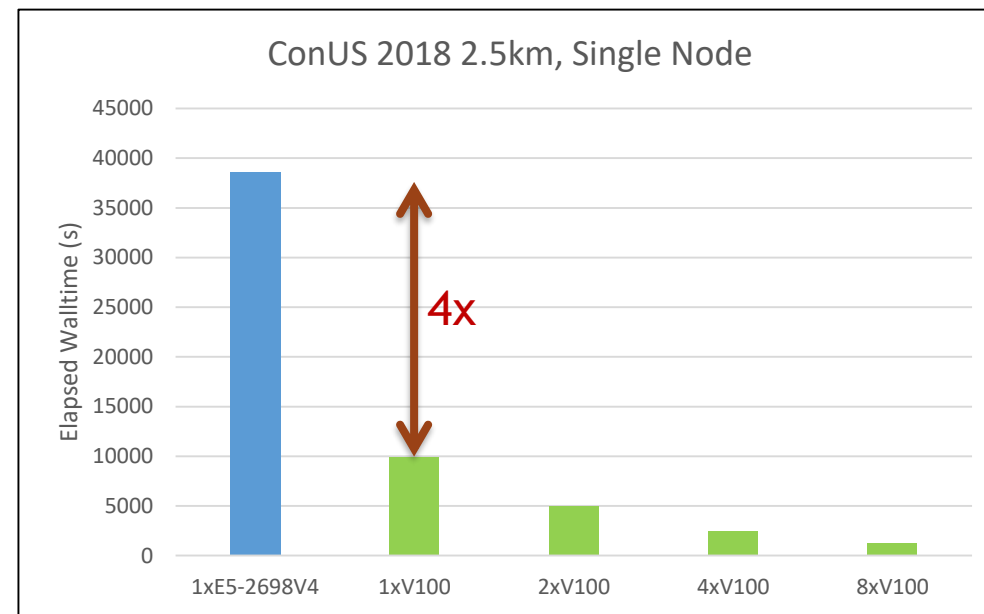
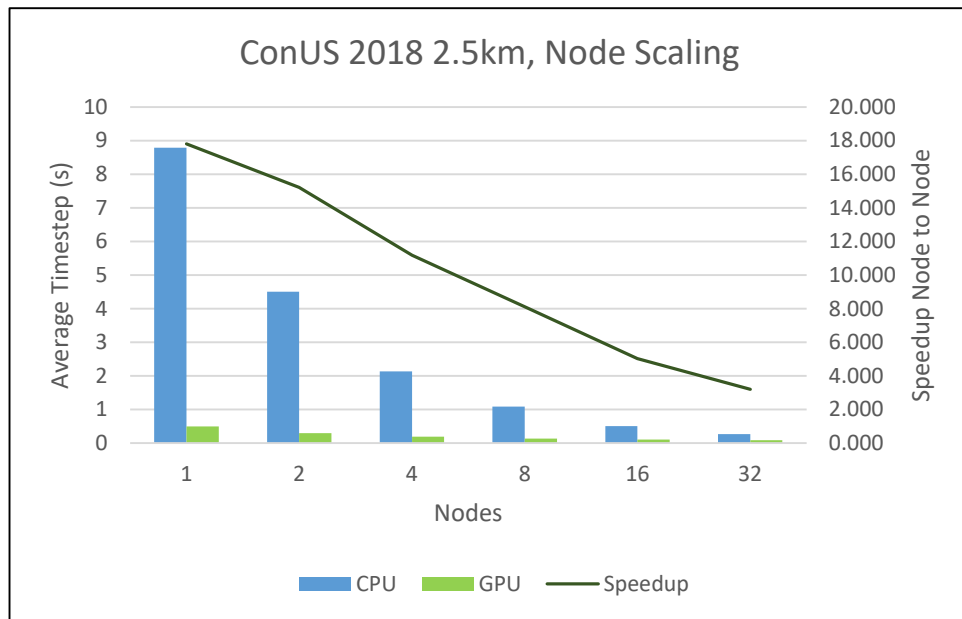
5-layer TDS	1
-------------	---

Cumulus

WRFg Full Model Performance



Smaller is better



1x CPU: 10hr 42m, 1x GPU: 2hr 45m, 4x GPU: 41m, 8x GPU: < 21 m

ConUS 2.5km 2018* Model, 2xE5-2698V4 CPU, 8xVolta V100 GPU

*Modifications to standard ConUS 2.5: WSM6, RRTMG, radt=3, 5-layer TDS

TOPICS OF DISCUSSION

- NVIDIA HPC UPDATE
- GPU RESULTS FOR NUMERICAL MODELS: IFS, WRF
- **GPU BENEFITS FOR AI MODELS**
David Hall, Ph.D.,
Sr. Solutions Architect, NVIDIA, Boulder, CO, USA