
HotGPU: A Thermal Profile Dataset for Immersion-Cooling AI Datacenters

Hengjia Zhang, Shuntao Zhu, Rui Lu, **Dan Wang**

The Hong Kong Polytechnic University

The Hong Kong University of Science and Technology

Growing Power Density in AI Datacenters

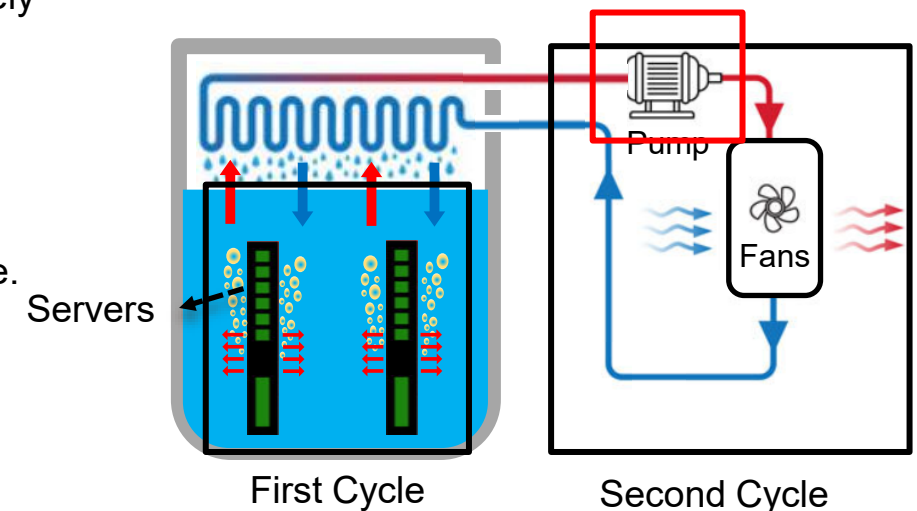
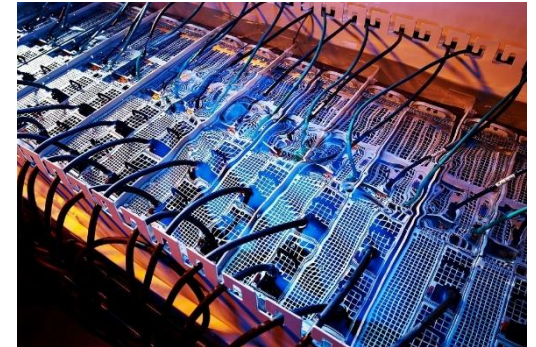
- The power density of AI servers increased **11x** from 2020 to 2025 [1].
 - CPU is 100W and GPU nowadays is 800W.
- High GPU power density creates thermal management challenges
 - Energy: Cooling has become a key component of datacenter operational cost.
 - Performance: Performance can decrease, e.g., thermal throttling in high temperature, etc. The cooling capacity can no-longer be assumed unlimited.

[1] IEA, Key Questions on Energy and AI

[2] A Thermal-aware Workload Scheduler for High-performance LLM Inference in Cooling-regulated Datacenters, HotCarbon'25

The new generation: Two-phase Immersion Cooling

- Servers/GPU racks are **immersed in dielectric liquid (Coolant)**: Heat from hardware is absorbed by the coolant
- Two-phase immersion cooling system in a nut shell:
 - First Cycle
 - GPU computing (e.g., LLM training) generate heat.
 - Heat absorbed by coolant and **phase change** from liquid to vapor (extremely efficient in heat absorption).
 - Vapor touches condenser pipe.
 - Vapor condenses back to liquid.
 - Second Cycle
 - Heat transferred from vaporized coolant to the water in the condenser pipe.
 - Hot water circulates through the condenser, driven by a pump.
 - Heat from the water removed to the ambient air by fans



To remove 2000W heat [1][2]:

Immersion cooling	vs.	Air cooling
59W	vs.	1020W

[1] Eaton. 2024. AC Unit for Server Racks - Rack Mount, 7,000 BTU (2.0 kW), 120V, 8U. <https://assets.tripplite.com/product-pdfs/en/srcool7krm.pdf>.

[2] Cheng Liu and Hang Yu. 2021. Evaluation and optimization of a two-phase liquid-immersion cooling system for data centers. *Energies* 14, 5 (2021), 1395. 3

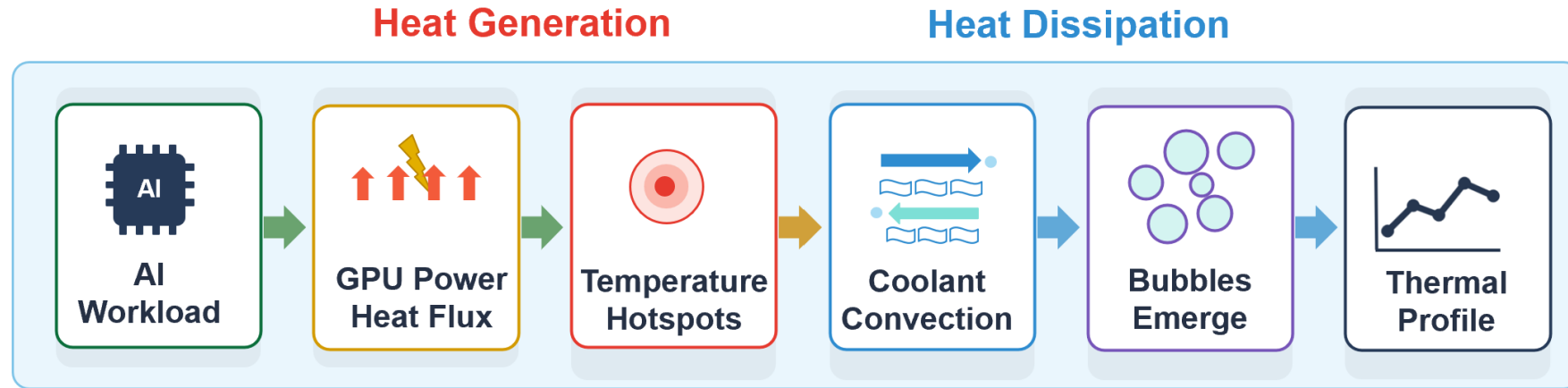
Challenge for immersion cooling AI datacenter research

- ❑ Evaluation is difficult
- ❑ Real-world experiments:
 - ❑ Expensive
 - ❑ GPU vendors may stop warranty if their GPU is immersed
- ❑ Data-driven: Limitation in existing datasets
 - ❑ AI workload/GPU power datasets: No immersion cooling data
 - ❑ Immersion cooling datasets: Often use constant power to generate heat

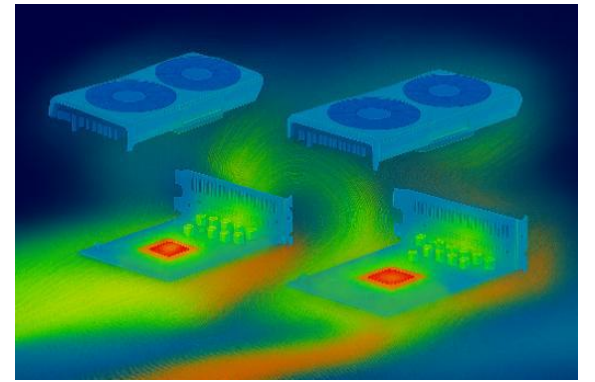
There is a gap in datasets and this is the target of this paper

To generate a dataset through CFD

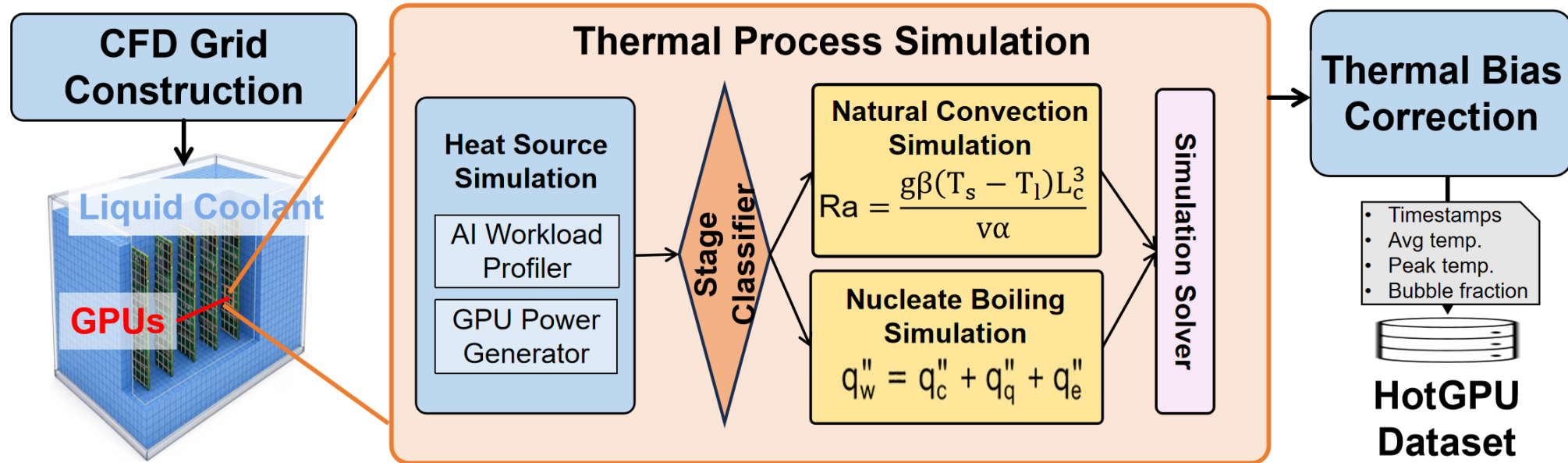
The cooling is a thermal process with heat generation and heat dissipation



We simulate the thermal process by computational fluid dynamics (CFD)



HotGPU: Overview



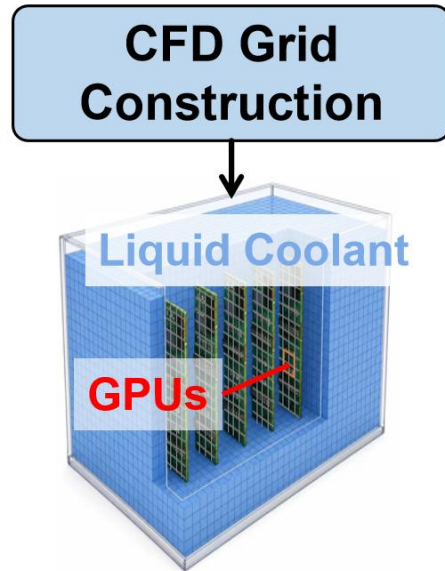
CFD is computation expensive. It took 3742-hour in a high performance 128 core CPU.

HotGPU: Overview

- ❑ Time (10ms)
- ❑ Power
- ❑ Average temperature (°C)
- ❑ Peak temperature (°C)
- ❑ Fraction of bubbles
- ❑ Kernel types (GMM, ReLU, Softmax, etc.)

The data is the GPU temperature. We didn't maintain the data for each cell since scheduling will be applied to GPU as a whole. The cells are used for simulation (e.g., heat convection).

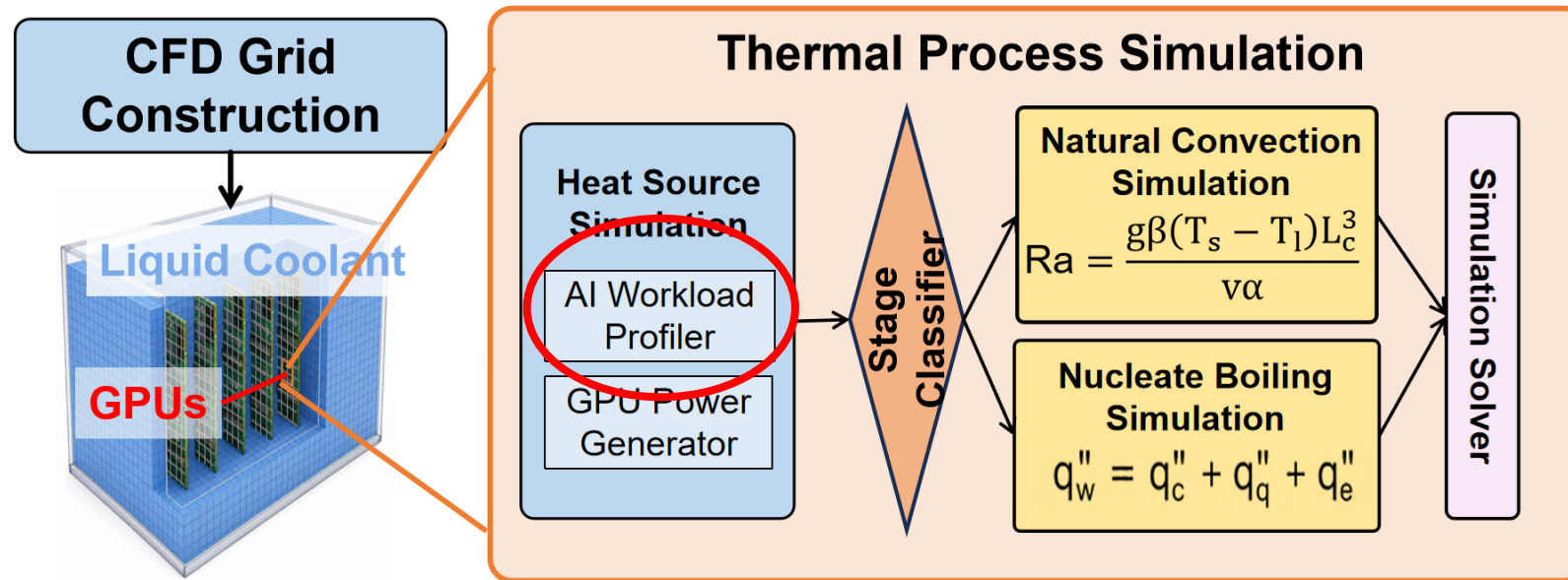
HotGPU: Grid Construction



■ CFD Grid Construction

- To construct the grid given a 3D geometric environment of an immersion cooling tank
- To configure the shape, size of the tank, GPUs, coolant region, and condenser pipe.
- To find the resolution of the cells of the grid: **0.5mm** at GPU and **15mm** at the periphery
- The state-of-the-art software will generate the grid structure.

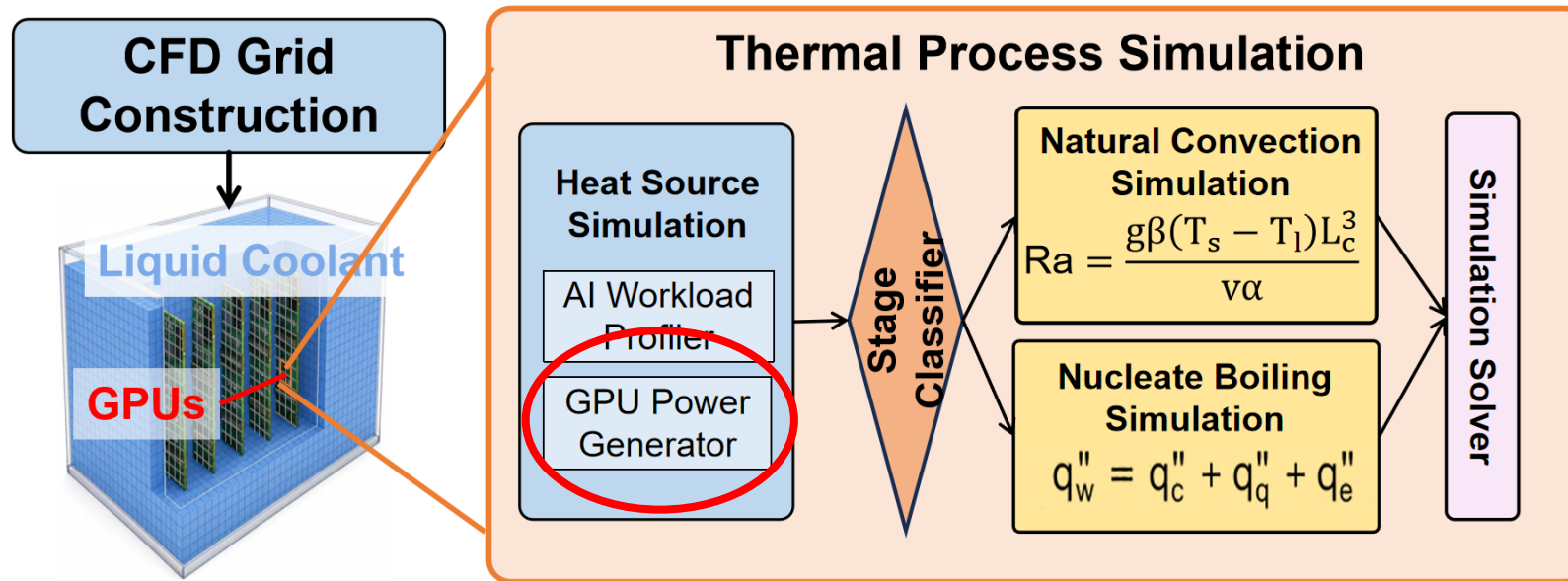
HotGPU: Thermal Process Simulation



■ Thermal Process Simulation

- To develop a mapping of the execution of an AI model to a time-series of execution of kernel functions (GMM, ReLU, Softmax, etc.) by Pytorch profiler

HotGPU: Thermal Process Simulation

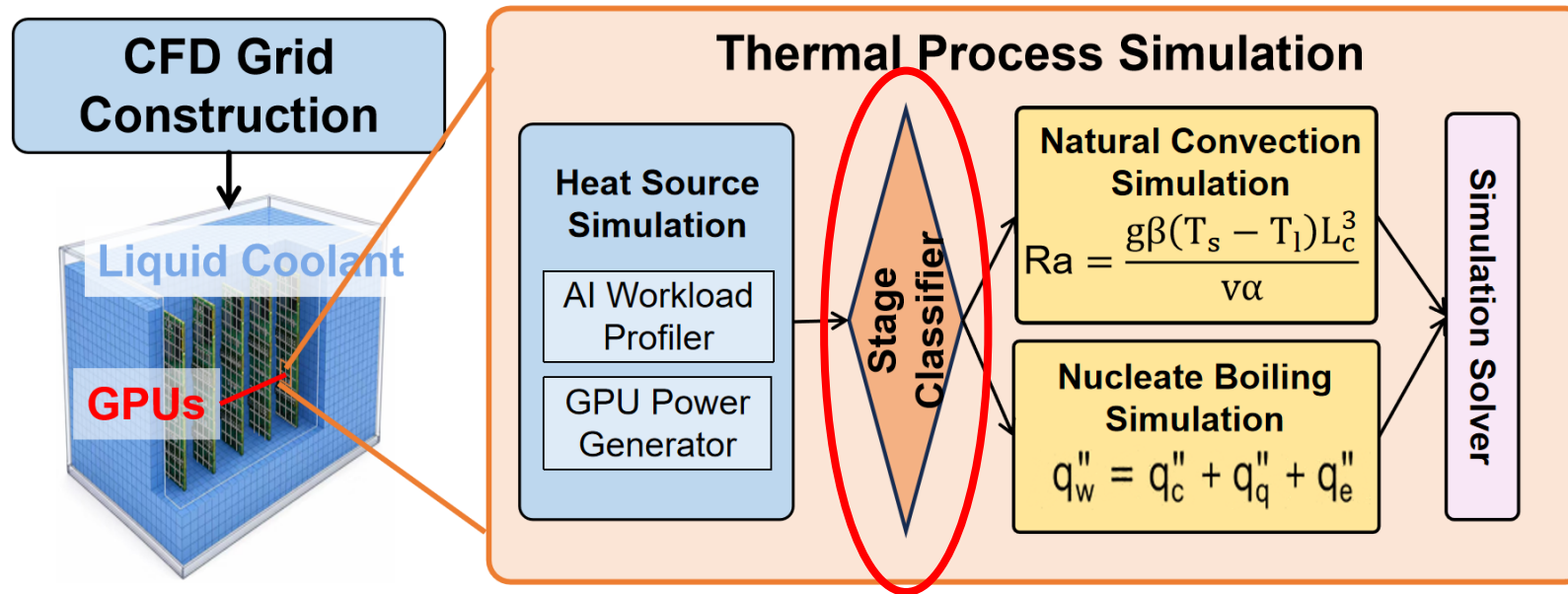


■ Thermal Process Simulation

- (Power to kernel) Given a power trace (which is most commonly available), to develop a mapping of the power to kernel functions through a handcraft power feature extraction and mapping.
- (Power to heat) Given a power trace, to construct the (volumetric) heat flux of a GPU by

$$q''' = P(t)/V_{\text{die}},$$

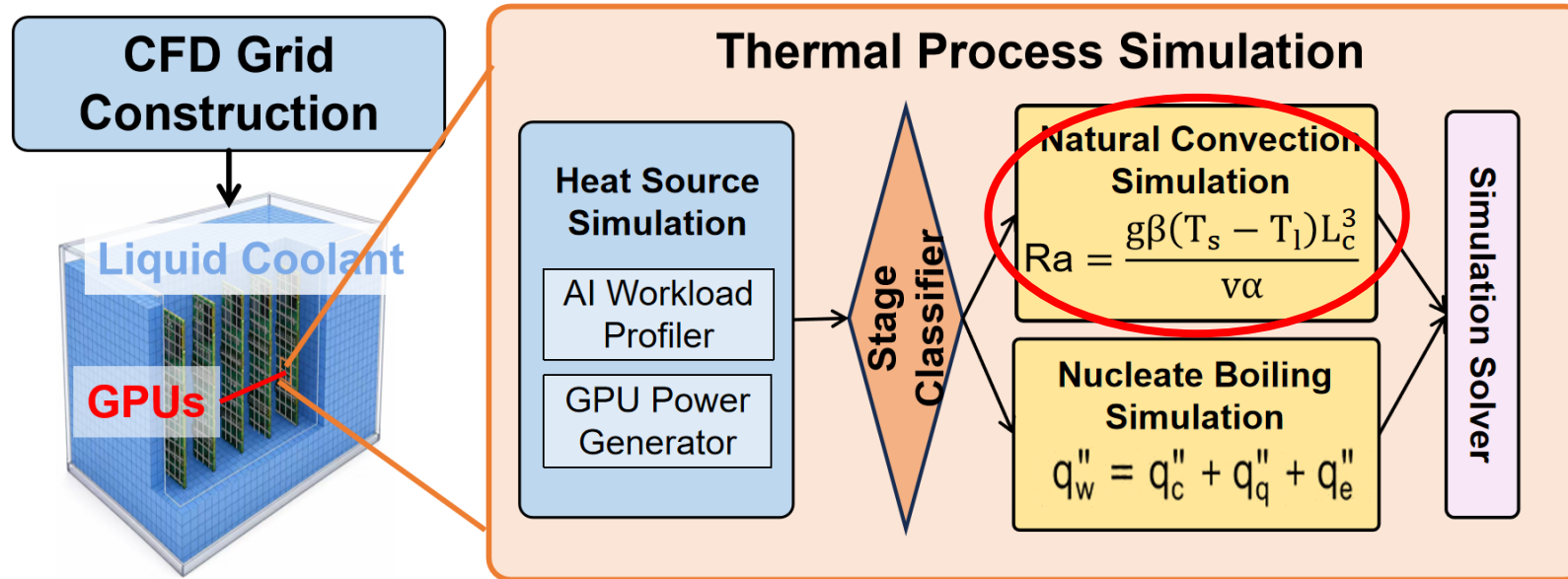
HotGPU: Thermal Process Simulation



■ Thermal Process Simulation

- Calculate the heat dissipation stage of the immersion cooling by the classification threshold of the GPU temperature

HotGPU: Thermal Process Simulation

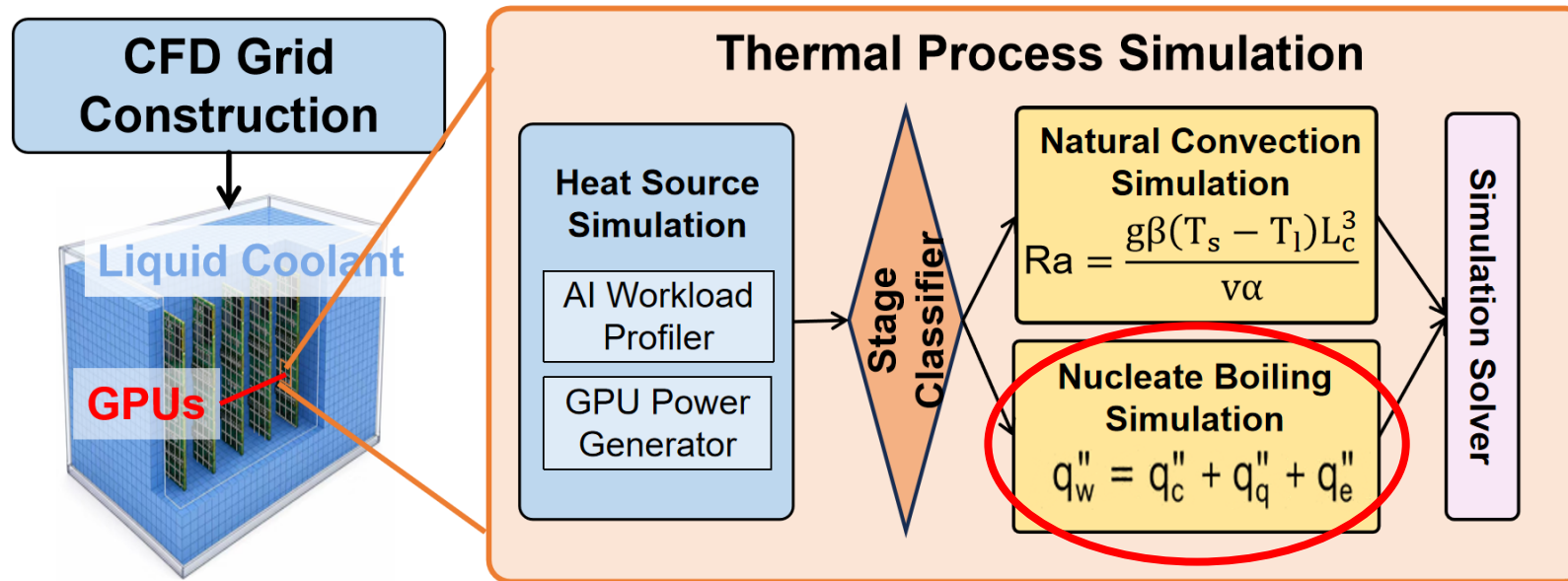


■ Thermal Process Simulation

- Simulate the natural convection process in the heat dissipation of immersion cooling by physics, e.g., a representative one is

$$Ra = g\beta(T_s - T_l)L_c^3 / (v\alpha),$$

HotGPU: Thermal Process Simulation

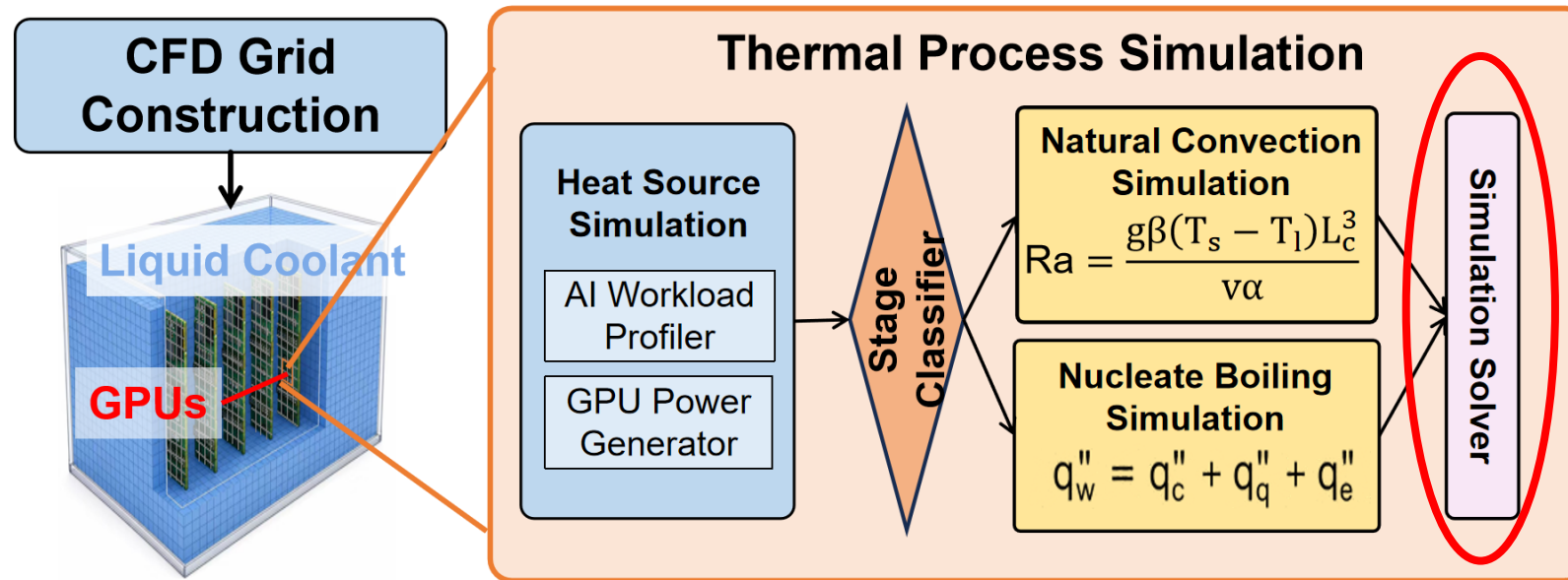


■ Thermal Process Simulation

- Simulate the nucleate boiling process in the heat dissipation of immersion cooling by physics, e.g., a representative one is

$$q_w'' = q_c'' + q_q'' + q_e'',$$

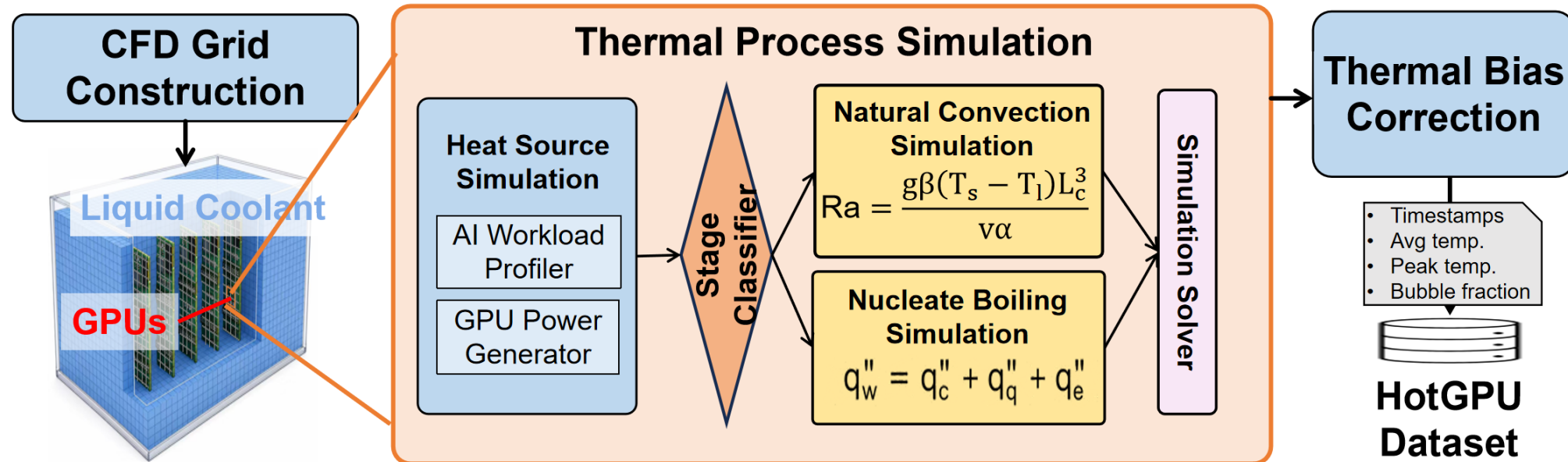
HotGPU: Thermal Process Simulation



■ Thermal Process Simulation

- To update thermal and flow fields through the equations of energy, momentum, mass conservation over all CFD cells – this is assisted by state-of-the-art CFD software

HotGPU: Thermal Bias Correction



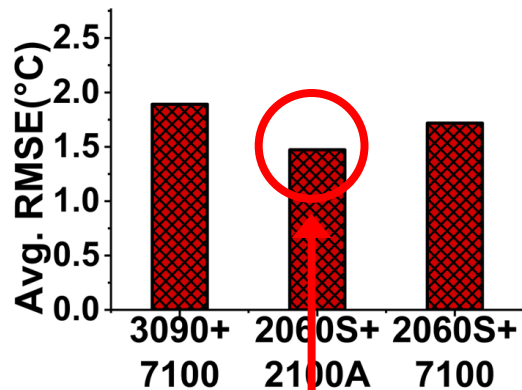
■ Thermal Bias Correction

- Correct GPU-specific temperature bias caused by simplified CFD modeling.
- A data-driven calibration.

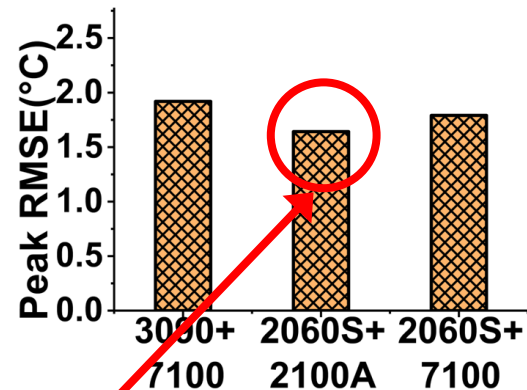
Validation with Real-World Measurements

■ Testbed:

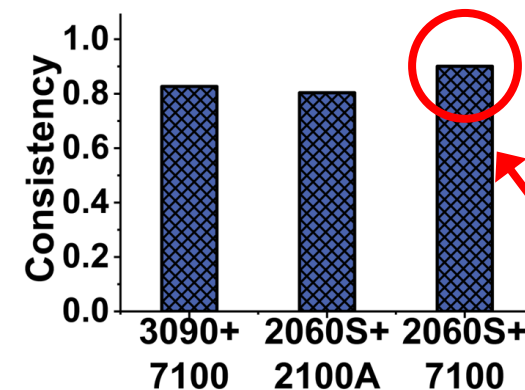
- **GPU-Coolant cases:** (1) RTX 3090+Novec-7100, (2) RTX 2060S+Novec-7100, (3) RTX 2060S + Noah-2100A.
- **Metrics:** RMSE, Thermal transition consistency: whether turning points are captured, i.e., the temperature shifts from increasing to decreasing, or vice versa



(a) Avg. RMSE.



(b) Peak RMSE.



(c) Consistency.

An RMSE of **1.47°C** and **1.64°C** for the average and peak temperature. HotGPU achieves **90.06%** on thermal transition consistency.

Case Study 1 - Capacity Planning

- Objective: To choose server spacing.
 - Thermal headroom: the temperature for safe AI computing (above which, there is low impact of neighboring GPUs)
- Fig. 1 shows the temperature as a function of spacing for RTX 2060S+Noah-2100A.
- Results:
 - HotGPU reveals how spacing affects average and peak temperature of a GPU (where there is a symmetric neighboring GPU)
 - Beyond 40mm, further spacing brings marginal temperature difference

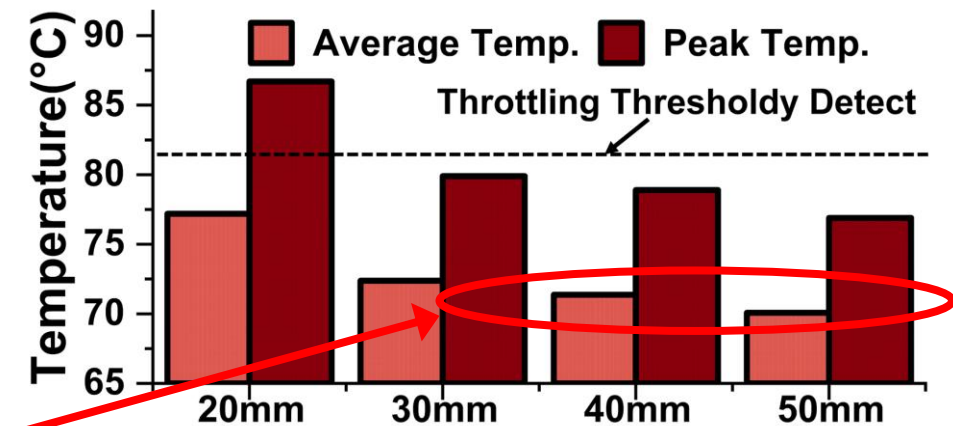


Fig. 1

Case Study 2 - Thermal-Aware Scheduling

- Objective: To detect thermal throttling risks.
- Fig. 2 shows temperature of an AI workload (a training trace of a diffusion model) as a function of time for H100+Novec-7100
- Results:
 - Without HotGPU, missing or delayed throttling
 - With HotGPU, throttling detected

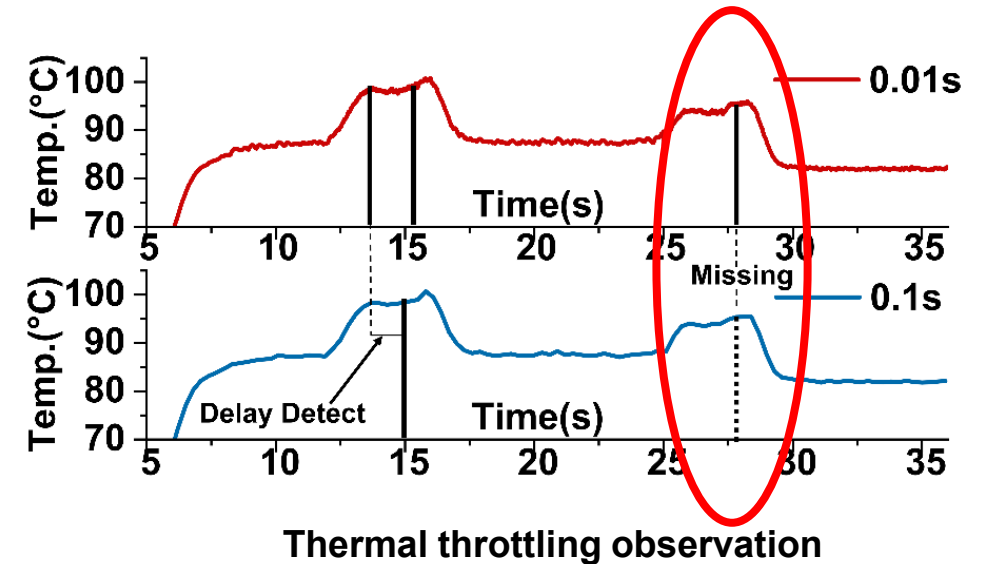


Fig. 2

Conclusion

- We develop HotGPU, an end-to-end thermal profile dataset for immersion-cooled AI datacenters.
- HotGPU uses high-fidelity CFD simulations validated by real-world experiments.
- HotGPU can assist capacity planning and thermal-aware scheduling.

Thank you!
Q&A