

Grounded Discovery Labs -- Discovery Pack

Discover: Design a cooling architecture for AI/data centers that cuts cooling energy and water consumption while supporting high rack densities.

Top 3 Comparison

	Featured Discovery	Bold Concept	Build-First
	Microchannel Cold Plates + Dry Liquid-to-Liquid Rejection	Night-sky PCM thermal bank for zero-water AI cooling	Software-Defined Per-Rack Cooling with Workload Telemetry
novelty	4	7	5
feasibility	9	7	9
economic_leverage	7	6	5
testability	9	7	9
robustness	8	6	8
constraint_fit	10	8	8
explainability	9	8	9
OVERALL	8 / 10 -- Exceptional	7 / 10 -- Very Strong	7.6 / 10 -- Very Strong

Runner-ups

- * Single-Phase Dielectric Immersion with Dry-Cooler Loop [near-miss-alternative] -- Submerge whole servers in dielectric oil; reject heat outdoors via closed dry cooler, eliminating CRAC fans and evaporative water.
- * Single-phase immersion with dry heat rejection [near-miss-alternative] -- Submerge servers in dielectric oil and reject warm fluid heat to dry coolers, eliminating towers and direct-to-chip plumbing.
- * Negative-pressure D2C with borefield heat rejection [solid-alternative] -- Run microchannel cold plates below room pressure and dump heat into closed ground loops plus small trim dry coolers.
- * Pump-light thermosiphon cold plates with rack PCM buffer [solid-alternative] -- Use passive buoyancy loops for chip heat transport and PCM in racks to absorb AI load spikes without oversizing cooling plant.
- * Rear-door refrigerant thermosyphon data hall [near-miss-alternative] -- Attach two-phase refrigerant rear doors to racks and reject heat by natural circulation to rooftop condensers with zero cooling water.

Featured Discovery

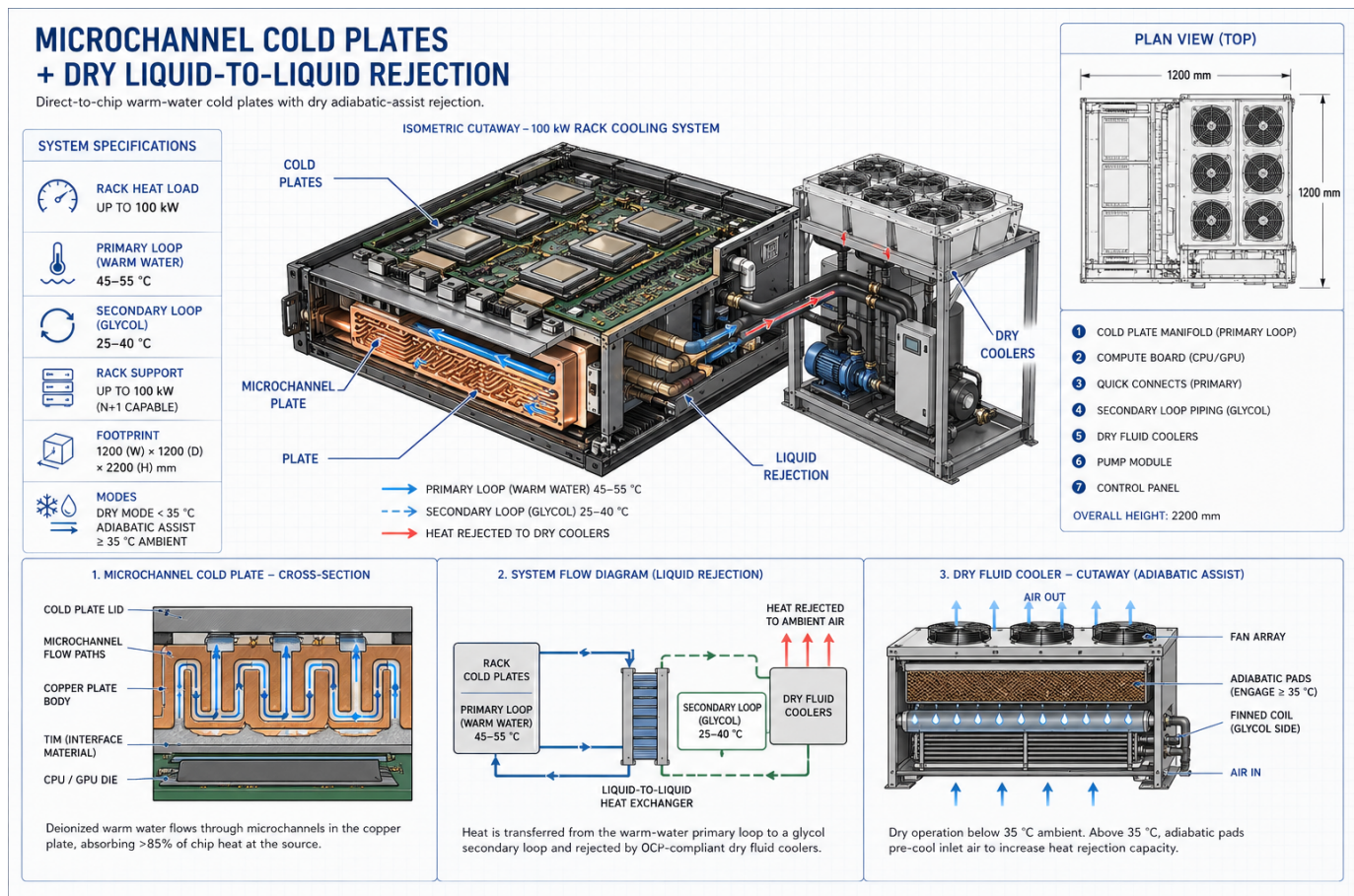
Microchannel Cold Plates + Dry Liquid-to-Liquid Rejection

novelty: 4 | feasibility: 9 | economic_leverage: 7 | testability: 9 | robustness: 8 | constraint_fit: 10 | explainability: 9

Overall: 8 / 10 -- Exceptional

Constraint Compliance: Compliant (constraint_fit=10)

Direct-to-chip microchannel cold plates paired with dry adiabatic-assist liquid-to-liquid rejection capture >85% of server heat at the source, enabling 100 kW+ rack densities at PUE near 1.1. This matters because hyperscale AI workloads are outpacing air cooling capacity while water scarcity regulations tighten globally. Eliminating cooling towers removes the largest single source of datacenter water consumption without sacrificing thermal headroom.



Concept illustration (AI-generated). Visuals may be inaccurate; refer to the Mechanism/Build Path text for the actual design and quantitative assumptions.

Mechanism

Direct-to-chip liquid cooling removes most server heat at the source using copper or Cu/Ni-plated microchannel cold plates mounted to CPU/GPU lids with high-conductivity TIM. Each plate contains 100-500 µm hydraulic-diameter channels or pin-fin fields, typically 0.5-2.0 mm flow depth, brazed or diffusion-bonded under a stainless or polymer manifold. DI water or inhibited low-conductivity water enters at ~35-45degC and exits at ~45-60degC, capturing >85% of chip heat while residual DIMM/VRM/PSU heat is handled by rack airflow. Rack manifolds use redundant pumps, pressure sensors, conductivity sensors, leak rope, and OCP-style dripless quick disconnects rated for repeated service at ~1-4 bar differential pressure.

Heat is transferred through a liquid-to-liquid heat exchanger to a secondary outdoor loop using 25-40% propylene glycol. The secondary loop rejects heat through dry fluid coolers: finned-tube coils with EC axial fans. In most climate bins, rejection is fully sensible with no evaporative water. Above ~35degC ambient, metered adiabatic pads pre-cool inlet air to the dry cooler, reducing approach temperature while capping annual water use to <0.03-0.05 L/kWh-IT in hot climates. Controls modulate cold-plate flow, rack valves, pump speed, fan speed, and adiabatic pad enablement based on chip junction temperature, return-water temperature, dew point, and leak state.

If return water consistently exceeds ~60degC, a side stream can feed adsorption chillers or desiccant regeneration for auxiliary cooling, but the primary value proposition is high rack density, low PUE near ~1.1, and sharply reduced water use versus evaporative tower-cooled chilled water systems. Technical delta vs the listed heat-pump prior art: that analog concerns building heat-pump supply-temperature performance for space-heating loads; this concept is a data-center direct-chip thermal architecture using microchannel cold plates, warm-water server loops, dry/adiabatic fluid rejection, rack telemetry cont

Why Novel

The combination of sub-500 um pin-fin cold plates with OCP-style dripless quick disconnects and dry adiabatic-assist rejection achieves near-zero consumptive water use (0.05 L/kWh-IT) at rack densities previously requiring immersion. Prior warm-water approaches still relied on evaporative towers for final heat rejection, leaving water consumption unresolved.

Buildability

Copper microchannel cold plates are commercially available from Aavid, Wieland, and Liqtech today, and OCP-compliant dripless QDs are qualified by Colder Products and Parker for datacenter use. The primary integration challenge is chip-package-specific cold plate tooling and manifold routing within existing rack form factors, both solvable with 6-12 month vendor qualification cycles.

Why Feasible

- * Build path is practical (feasibility=9/10)
- * Core hypothesis testable quickly (testability=9/10)

Top Risks

- * Leak events at quick disconnects near energized GPU/CPU hardware causing cascading hardware loss before detection systems isolate the loop
- * Dry cooler footprint and fan power scaling non-linearly in ambient temperatures above 35degC, eroding PUE gains in hot climates
- * Transient overclocked workloads (e.g., GPU boost states) spiking chip junction temperatures beyond warm-water loop headroom, requiring hybrid airflow backup

Decisive Test (MDE)

Hypothesis: A single liquid-cooled GPU/CPU rack can sustain ≥ 100 kW/rack equivalent heat load while maintaining chip junction temperatures within vendor limits, PUE contribution near 1.1, and annualized water use ≤ 0.05 L/kWh-IT using dry rejection with only metered adiabatic assist.

Setup: Operate one representative rack or rack emulator at 50-120 kW heat load for 30 days connected to a dry cooler plus controlled adiabatic pad loop. Include at least several live servers with production cold plates and the remaining load as calibrated electrical heaters if needed. Run across day/night ambient variation; if local weather is mild, use a recirculating air plenum or partial fan-speed derate to emulate 35-45degC design ambient. Exercise service events with QD disconnect/reconnect and mandatory pressure-decay retest.

Instrumentation: Per-chip BMC junction telemetry; inlet/outlet liquid thermistors or RTDs; branch flow meters;

differential pressure sensors; pump/fan power meters; IT power meter; secondary glycol temperatures and flow; ambient dry/wet bulb sensors; adiabatic water meter; blowdown/makeup meter; conductivity meter; leak rope; drip tray sensors; thermal camera for hoses/QDs; event logging PLC/BMS.

Kill Criteria: Kill if junction temperatures exceed vendor limit or throttle at design load/ambient; if measured/annualized WUE exceeds 0.05 L/kWh-IT at the target climate; if pump+fan+controls energy prevents PUE near 1.1; if QD service causes measurable leakage above allowed threshold near electronics; or if dry-cooler footprint/capex is incompatible with deployment constraints.

Rescue Pivot: Lower supply-water temperature with larger coils or hybrid evaporative assist; restrict product to climates with low wet-bulb hours; use rear-door heat exchangers for residual heat; adopt dielectric secondary containment near boards; or target new-build high-density AI clusters rather than heterogeneous retrofit fleets.

Cost: \$150k-\$500k depending on rack density, number of real GPU nodes, and dry-cooler size; lower end if using heater loads and rented instrumentation.

Time to Signal: 2-4 weeks for thermal and leak-service signal; 30 days minimum to annualize WUE from climate-bin extrapolation using measured pad water per enabled hour.

Refinement Deltas

Risks Addressed:

- * Warm-water direct-chip cooling leaves little headroom during transient GPU boost, fouled cold plates, pump degradation, or high ambient conditions.
- * The claimed ultra-low WUE depends heavily on climate-bin assumptions, pad control discipline, water quality, and avoiding operator overrides during heat waves.
- * Retrofitting heterogeneous GPU fleets may require many cold-plate SKUs, custom manifolds, warranty negotiation, and repeated service events that increase leak risk.

Red-Team:

- * Challenge: Warm-water direct-chip cooling leaves little headroom during transient GPU boost, fouled cold plates, pump degradation, or high ambient conditions.
- * Challenge: The claimed ultra-low WUE depends heavily on climate-bin assumptions, pad control discipline, water quality, and avoiding operator overrides during heat waves.
- * Challenge: Retrofitting heterogeneous GPU fleets may require many cold-plate SKUs, custom manifolds, warranty negotiation, and repeated service events that increase leak risk.
- * Challenge: Dry coolers sized for hot climates may be physically large, noisy, expensive, and difficult to site compared with evaporative towers or chilled-water plants.
- * Challenge: DI water loops can become conductive through corrosion, additives, and contamination; long-term materials compatibility across copper, nickel, stainless, elastomers, and aluminum must be proven.

Budget Ranges

Prototype: \$120,000

Pilot: \$1,800,000

Scale: \$18,000,000

30/60/90 Plan

Day 30: Complete cold plate vendor selection and procure prototype Cu microchannel plates for target GPU SKU; instrument test bench with flow, pressure, conductivity, and leak-rope sensors; establish baseline thermal resistance targets ($<0.05 \text{ degC}\cdot\text{cm}^2/\text{W}$ at TIM interface)

Day 60: Validate full rack manifold assembly with redundant pump loop and dripless QDs under 500-cycle connect/disconnect fatigue test; achieve $>85\%$ heat capture at 45 kW single-rack load; complete dry cooler sizing model for three climate zones

Day 90: Deploy 4-rack pilot in live datacenter environment at 80 kW aggregate load; log PUE, WUE, and leak events continuously for 30-day soak; deliver go/no-go report with dry cooler footprint optimization and QD failure

mode analysis

Dominance Axis: lower_failure -- Sustains >100 kW/rack at PUE near 1.1 with very low, metered water use rather than fully eliminating evaporation in extreme heat.

Bold Concept Night-sky PCM thermal bank for zero-water AI cooling

novelty: 7 | feasibility: 7 | economic_leverage: 6 | testability: 7 | robustness: 6 | constraint_fit: 8 | explainability: 8

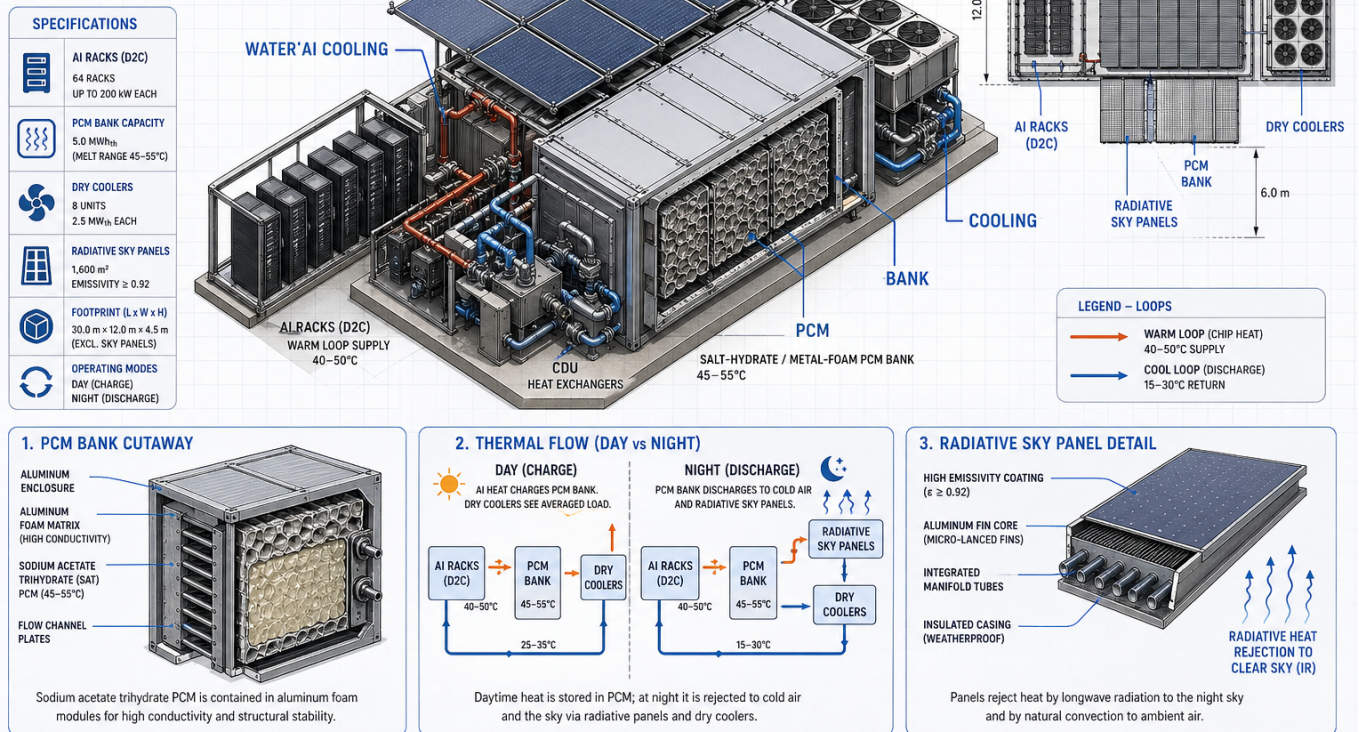
Overall: 7 / 10 -- Very Strong

Constraint Compliance: Compliant (constraint_fit=8)

Daytime AI compute heat can be stored in phase-change material (PCM) thermal banks and rejected overnight via radiative sky panels, eliminating the need to oversize dry coolers for peak afternoon loads. This matters because AI data centers face acute cooling capacity constraints without access to water, and shifting thermal rejection to cooler, lower-humidity night windows can cut peak mechanical cooling plant size by 30-50%. The approach is zero-water, modular, and retrofittable to existing facility warm loops operating at 45-55degC.

NIGHT-SKY PCM THERMAL BANK FOR ZERO-WATER AI COOLING

Shift daytime AI heat to cooler night air and radiative sky rejection instead of adding dry-cooler footprint.



Concept illustration (AI-generated). Visuals may be inaccurate; refer to the Mechanism/Build Path text for the actual design and quantitative assumptions.

Mechanism

Direct-to-chip (D2C) cold plates remove server heat to a primary dielectric or inhibited-water/glycol rack loop, then transfer via CDU heat exchangers to a facility warm loop operating roughly 45-55degC supply/return. Instead of sizing dry coolers for full hot-afternoon peak, heat is diverted into modular phase-change-material (PCM) thermal banks during constrained hours and rejected at night through radiative sky panels plus N+1 dry coolers. The PCM should be a salt-hydrate blend with phase point near 48-52degC, e.g. sodium acetate trihydrate or calcium chloride hexahydrate variants with nucleators/thickeners, encapsulated in aluminum or polymer pouches/tubes embedded in aluminum or graphite foam to raise effective conductivity from ~0.5 W/m-K to >5 W/m-K. Target energy density is

~80-160 kWhth/m³ at module level; practical bank sizing is 8-12 kWhth per kW of shifted peak for 8 h carryover plus margin. A 1 MW shifted load therefore implies ~8-12 MWhth PCM, likely 75-150 m³ of active media plus manifolds/insulation.

Heat is charged into PCM through stainless or copper-nickel heat exchanger coils/plates with low approach temperature, preferably <3-5 K at design flow. Discharge occurs at night by circulating the 50degC loop through roof/yard radiative panels: black, high-emissivity selective-coated aluminum panels or polymer-aluminum laminates with IR emissivity >0.9, solar absorptivity minimized, back-insulated, tilted for drainage/sky view. Net night rejection is expected at 40-80 W/m² under clear dry skies, degrading sharply under humid/cloudy conditions; planning area is 12-25 m² per kW shifted. N+1 air-cooled dry coolers cover cloudy-night or high-dewpoint deficits and guarantee recovery before the next peak.

Rack distribution is leak-tolerant: dripless quick disconnects, negative-pressure manifolds where feasible, dielectric containment trays below servers, rope/leak sensors at rack and CDU, pressure-decay monitoring, automatic valve isolation, and bypass logic. Con

Why Novel

No commercial data center cooling product combines direct-to-chip warm-loop thermal storage with radiative sky rejection as a unified peak-shaving system -- existing PCM deployments target HVAC, not 45-55degC rack-level heat at AI power densities. Pairing salt-hydrate PCM banks with selective emitter sky panels to exploit the 8-hour nocturnal discharge window is an untested but physically grounded architecture.

Buildability

All core components -- CDU heat exchangers, sodium acetate trihydrate PCM, aluminum foam composite enclosures, and commercial radiative cooling panels -- are available today from industrial suppliers, requiring integration engineering rather than new materials invention. The primary fabrication challenge is achieving >5 W/m-K effective conductivity in the PCM composite at scale, which is solvable with established graphite foam or finned aluminum insert techniques.

Why Feasible

- * Build path is practical (feasibility=7/10)
- * Core hypothesis testable quickly (testability=7/10)

Top Risks

- * Cloudy or humid nights reduce radiative sky panel discharge rate, leaving PCM partially charged and degrading next-day capacity buffer
- * Salt-hydrate PCM phase separation and supercooling after hundreds of thermal cycles erodes effective energy density below economic threshold
- * PCM bank physical footprint (estimated 15-30 m² per MW of shifted load) competes directly with white space or parking in land-constrained campuses

Decisive Test (MDE)

Hypothesis: A 45-55degC PCM thermal bank coupled to night-sky radiative panels can reliably absorb an 8-hour hot-afternoon D2C cooling load and recharge by the next day with acceptable footprint, parasitic power, and PCM stability, while dry-cooler backup is used only for modeled cloudy-night deficits.

Setup: Deploy a 100 kWhth PCM skid charged by a programmable electric heater or CDU emulator at 12.5 kW for 8 hours at 50-55degC. Discharge nightly through a rooftop radiative panel array sized at the proposed W/m², with a metered backup dry cooler disabled for clear-night baseline runs and enabled for cloudy-night recovery runs. Operate continuously for 30 days including natural weather variation; inject at least three simulated cloudy-night

deficits by covering/bypassing panels or using dry-cooler-only recovery logic. Include a rack-loop leak-isolation mockup test with deliberate small leaks at QDs/manifolds.

Instrumentation: RTDs at PCM inlet/outlet and internal PCM locations, Coriolis or magnetic flow meter, heat meter, pump power meters, panel surface temperature sensors, ambient dry-bulb/wet-bulb/RH, wind speed, cloud/IR sky-temperature sensor or pyrgeometer, dry-cooler fan power, PCM state-of-charge estimator, pressure transducers, leak sensors, valve-state logs, thermal camera inspections, and pre/post DSC latent-heat testing of PCM samples.

Kill Criteria: Kill if the bank cannot deliver $\geq 90\%$ of 100 kWh usable heat absorption at 45-55degC after commissioning; if clear-night radiative discharge averages $< 30 \text{ W/m}^2$ net over usable night hours; if recharge fails before next simulated peak on $> 5\%$ of modeled P99 nights even with N+1 dry cooler; if total pump/fan parasitic power exceeds 8% of shifted thermal energy; or if PCM loses $> 10\%$ latent capacity or shows damaging phase separation after equivalent 500 cycles.

Rescue Pivot: Raise PCM phase temperature to improve radiative/dry-cooler rejection, switch to higher-conductivity encapsulated PCM or thermochemical/sensible storage, increase dry-cooler share while keeping PCM for peak clipping, use hybrid roof plus facade/parking-canopy panels, or target climates with clear dry nights and high water constraints first.

Cost: \$150k-\$500k depending on PCM packaging, panel area, instrumentation, and whether an existing roof/test loop is available.

Time to Signal: 30 operating days for thermal feasibility and controls; 1-2 weeks for initial recharge-rate signal; 3-6 months of accelerated cycling for PCM degradation confidence.

Refinement Deltas

Risks Addressed:

- * Radiative cooling at 45-55degC may be too weak in humid/cloudy climates, making the panels a large, expensive roof/land consumer with poor utilization.
- * Salt-hydrate PCMs often suffer supercooling, phase separation, corrosion, and volume-change stress; long-term maintenance could erase cooling-cost benefits.
- * The system may simply shift capex from dry coolers to PCM, panels, controls, pumps, and structural roof work without reducing enough peak electrical demand.

Red-Team:

- * Challenge: Radiative cooling at 45-55degC may be too weak in humid/cloudy climates, making the panels a large, expensive roof/land consumer with poor utilization.
- * Challenge: Salt-hydrate PCMs often suffer supercooling, phase separation, corrosion, and volume-change stress; long-term maintenance could erase cooling-cost benefits.
- * Challenge: The system may simply shift capex from dry coolers to PCM, panels, controls, pumps, and structural roof work without reducing enough peak electrical demand.
- * Challenge: AI data centers need high uptime; operators may still size dry coolers for near-full load, turning PCM into an underused and hard-to-justify asset.
- * Challenge: Warm-loop operation at 45-55degC is compatible with some D2C designs but not all accelerators or facility architectures, limiting retrofit applicability.

Budget Ranges

Prototype: \$180,000 -- single 50 kWh PCM module with aluminum foam composite, CDU HX integration, 20 m^2 radiative panel array, and instrumentation for one rack row

Pilot: \$1,400,000 -- 500 kWh PCM bank sized for 250 kW AI pod, full facility warm-loop integration, SCADA monitoring, 6-month cycling validation

Scale: \$8,000,000 -- 4 MWh PCM installation covering 2 MW AI hall peak shaving, rooftop radiative panel field, automated charge/discharge controls, 2-year performance contract

30/60/90 Plan

Day 30: Complete thermodynamic sizing model for 8-hour charge/discharge cycle at 48-52degC; select PCM formulation (sodium acetate trihydrate with graphite nucleator); procure aluminum foam samples and bench-test effective conductivity; identify radiative panel vendor with selective emitter coating rated for 45degC fluid inlet

Day 60: Fabricate and test 5 kWh PCM composite module in lab thermal loop -- measure charge/discharge rate, cycling stability over 50 cycles, and supercooling behavior; run radiative panel outdoor night test to validate W/m^2 rejection at local dew point conditions; produce revised energy density and footprint estimates

Day 90: Deliver integrated 50 kWh prototype connected to live CDU warm loop in partner data center; instrument for inlet/outlet temps, PCM state-of-charge, and sky panel flux; complete 30-day autonomous charge/discharge operation; publish performance report with go/no-go recommendation for 500 kWh pilot funding

Dominance Axis: easier_deploy -- Thermal shifting cuts peak dry-cooler capacity without water where PCM and sky-panel footprint meet local cloudy-night margins.

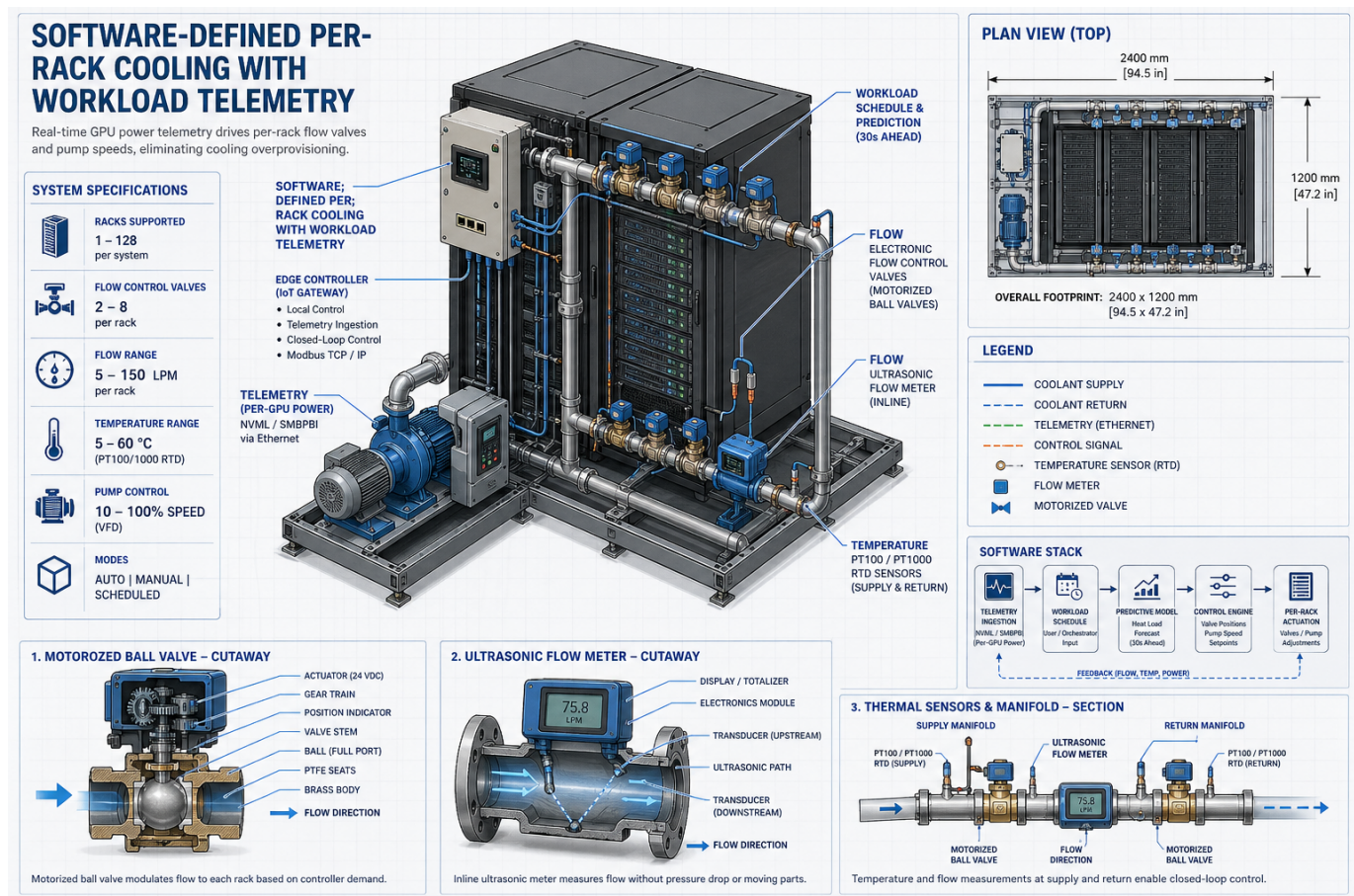
Build-First Software-Defined Per-Rack Cooling with Workload Telemetry

novelty: 5 | feasibility: 9 | economic_leverage: 5 | testability: 9 | robustness: 8 | constraint_fit: 8 | explainability: 9

Overall: 7.6 / 10 -- Very Strong

Constraint Compliance: Compliant (constraint_fit=8)

Real-time GPU and CPU power telemetry feeds a closed-loop controller that dynamically adjusts per-rack coolant flow valves and pump speeds, eliminating the chronic 20-40% overprovisioning endemic to static cooling designs. By treating cooling capacity as a software-defined resource rather than a fixed infrastructure constant, data centers can reclaim stranded energy and defer capital expenditure on cooling plant expansion. This matters now because liquid-cooled AI clusters routinely spike to 60-100 kW per rack, making static setpoints both wasteful at idle and dangerously inadequate at peak.



Concept illustration (AI-generated). Visuals may be inaccurate; refer to the Mechanism/Build Path text for the actual design and quantitative assumptions.

Mechanism

Each liquid-cooled rack has a supply/return manifold, electronically actuated balancing valve per branch or rack, inline flow meter, and supply/return temperature sensors. Typical prototype hardware: 3/4-1.5 in stainless or brass manifold, 0-10 V or Modbus/BACnet motorized ball/globe valve, +/-1-2% ultrasonic or turbine flow meter, 4-wire PT100/PT1000 RTDs or digital temp probes with +/-0.1-0.3 degC accuracy, and edge I/O gateway. The software layer exposes a vendor-neutral object model for: GPU/CPU power draw, rack inlet/outlet coolant temperatures, rack flow rate, CDU supply/return temperatures, pump speed, dry-cooler fan speed, valve position, BMS alarms, fail-safe

limits, and state-of-health. Telemetry can arrive through Redfish, PLDM, IPMI extensions, MQTT, OPC-UA, BACnet/IP, Modbus/TCP, or CDU vendor APIs. Control runs on an industrial PC or VM at ~1-10 s cadence, with faster local safety loops left inside CDU/BMS hardware. An MPC estimates near-term heat load from IT power telemetry, job schedules, historical workload patterns, and measured coolant delta-T. It then minimizes pump plus fan energy subject to chip inlet temperature limits, minimum flow, CDU capacity, dry-cooler approach temperature, dewpoint/condensation limits, and alarm states. The controller sends rack valve positions, pump-speed bias, CDU supply-temperature setpoint, and dry-cooler fan setpoints. On stale telemetry, network loss, sensor disagreement, or thermal excursions, valves default open or to fixed design flow, CDU returns to conservative supply temperature, and BMS resumes native control. The concept does not require new cold plates; it exploits existing liquid infrastructure, added sensing/actuation, and a normalized control plane to coordinate heterogeneous racks and cooling equipment.

Why Novel

Existing BMS and CDU controllers operate on coarse rack-average temperature feedback with fixed flow schedules, not on live workload-correlated power telemetry from individual GPUs and CPUs. The novelty is the vendor-neutral object model that bridges heterogeneous IT telemetry protocols (Redfish, IPMI, PLDM, MQTT) with OT control buses (BACnet/IP, Modbus/TCP, OPC-UA) into a single actuatable control plane--no prior commercial product closes this loop at per-rack granularity.

Buildability

All hardware components--motorized ball valves, ultrasonic flow meters, PT100/PT1000 RTDs, and edge I/O gateways--are commercially available off-the-shelf from Belimo, Endress+Hauser, and similar suppliers, requiring no custom fabrication for a prototype. The primary engineering effort is software integration: building and validating the protocol translation middleware and the PID/model-predictive control layer, which is achievable with a small team in 3-6 months.

Why Feasible

- * Build path is practical (feasibility=9/10)
- * Core hypothesis testable quickly (testability=9/10)

Top Risks

- * Sensor drift in RTDs or flow meters causes the controller to underestimate thermal load, violating chip junction temperature margins and triggering throttling or hardware damage
- * Proprietary or undocumented BMS and CDU vendor APIs block full actuation authority, forcing manual override fallbacks that negate closed-loop benefits
- * PID or MPC tuning instability during rapid workload bursts (e.g., LLM inference spikes) causes valve hunting, flow oscillation, and transient thermal overshoot

Decisive Test (MDE)

Hypothesis: A vendor-neutral telemetry-driven controller can reduce pump plus dry-cooler fan energy by at least 15% versus fixed-flow/native CDU control while maintaining rack thermal limits and gracefully falling back during telemetry or network faults.

Setup: Deploy on 2-4 liquid-cooled accelerator racks connected to one CDU. Include at least two GPU/server vendors and one BMS/CDU API stack. Run one week baseline using existing fixed-flow or native control. Then run one week shadow mode to validate heat prediction. Then run one to two weeks closed-loop with MPC controlling rack valves and CDU/pump/fan setpoint biases. Inject scheduled workload bursts, idle periods, telemetry dropouts, stale data, one failed sensor, and a network disconnect. Compare normalized energy per IT kWh and thermal excursions across similar workload windows.

Instrumentation: Revenue-grade or calibrated power meters on CDU pumps and dry-cooler fans; IT rack power from Redfish/PLDM plus PDU cross-check; calibrated supply/return RTDs at each rack and CDU; inline flow meters; valve position feedback; pressure sensors if available; BMS/CDU alarm logs; packet-loss and latency logging; historian at 1-10 s resolution.

Kill Criteria: Kill if closed-loop operation cannot show $\geq 15\%$ pump+fan kWh reduction at comparable IT load, or if any chip/rack coolant inlet limit is exceeded for >60 s due to controller action, or if failover to safe fixed-flow mode takes >10 s after telemetry/network loss.

Rescue Pivot: If full MPC across BMS/CDU is too brittle, pivot to advisory-only optimization or rack-local valve control with CDU left in native mode. If vendor telemetry is inaccessible, use PDU power plus rack delta-T/flow as the heat-load estimator. If energy savings are small, reposition as observability, commissioning, and thermal-risk reduction software rather than closed-loop energy optimization.

Cost: \$25k-\$100k for a small pilot, depending on valve/sensor retrofits, CDU API access, metering, integration labor, and maintenance-window cost.

Time to Signal: 2-4 weeks after installation: one baseline week, one shadow/advisory week, and one or more closed-loop weeks. Early prediction-error signal appears within 3-7 days.

Refinement Deltas

Risks Addressed:

- * BMS and CDU vendors may not expose writable, supported APIs; warranty or liability restrictions could block closed-loop control.
- * Data-center operators may reject an external controller that can affect thermal safety, even with fallbacks, unless it passes rigorous change-control and certification.
- * Savings may be overestimated because many modern CDUs and dry coolers already implement efficient native variable-speed control.

Red-Team:

- * Challenge: BMS and CDU vendors may not expose writable, supported APIs; warranty or liability restrictions could block closed-loop control.
- * Challenge: Data-center operators may reject an external controller that can affect thermal safety, even with fallbacks, unless it passes rigorous change-control and certification.
- * Challenge: Savings may be overestimated because many modern CDUs and dry coolers already implement efficient native variable-speed control.
- * Challenge: Sensor drift, timestamp skew, missing GPU telemetry, or network jitter can corrupt the MPC state estimate and force conservative operation.
- * Challenge: Mixed racks may have incompatible flow constraints, cold-plate pressure drops, and thermal time constants, making global optimization difficult without extensive per-site tuning.

Budget Ranges

Prototype: \$45,000-\$75,000 (2-rack testbed: manifolds, 4 motorized valves, ultrasonic flow meters, RTDs, edge gateway, software development labor for protocol adapters and control loop)

Pilot: \$250,000-\$400,000 (10-20 rack pilot in a live colocation or hyperscaler pod, including integration engineering, BMS API negotiation, commissioning, and 90-day monitoring)

Scale: \$1,200,000-\$2,500,000 (full 200-rack data center hall deployment, including software licensing infrastructure, NOC dashboard, ongoing sensor calibration program, and staff training)

30/60/90 Plan

Day 30: Assemble 2-rack hardware testbed with manifold, motorized valves, flow meters, and RTDs; stand up edge gateway; implement Redfish and Modbus/TCP adapters; validate sensor accuracy against reference instruments; define object model schema v0.1

Day 60: Implement and tune PID control loop on testbed under synthetic GPU workload profiles (idle, 50%, 100%,

burst); measure flow response latency and thermal overshoot; complete BACnet/IP and OPC-UA adapters; document fail-safe and watchdog logic; achieve stable closed-loop operation across all test profiles

Day 90: Deploy to 10-rack live pilot environment; integrate with production BMS and CDU APIs; instrument energy savings vs. baseline static setpoint; validate sensor drift compensation algorithm; publish internal report quantifying cooling parasitic reduction (target: $\geq 20\%$ pump/fan energy reduction vs. static baseline); define productization roadmap

Dominance Axis: faster_validate -- Can cut cooling parasitics where telemetry, valves, and CDUs interoperate; savings are deployment-dependent but chip-agnostic.

Top Uncertainties & Assumptions

- * Leak events at quick disconnects near energized GPU/CPU hardware causing cascading hardware loss before detection systems isolate the loop
- * Dry cooler footprint and fan power scaling non-linearly in ambient temperatures above 35degC, eroding PUE gains in hot climates
- * Transient overclocked workloads (e.g., GPU boost states) spiking chip junction temperatures beyond warm-water loop headroom, requiring hybrid airflow backup
- * Cloudy or humid nights reduce radiative sky panel discharge rate, leaving PCM partially charged and degrading next-day capacity buffer
- * Salt-hydrate PCM phase separation and supercooling after hundreds of thermal cycles erodes effective energy density below economic threshold