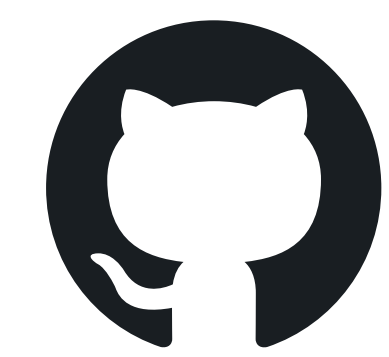


Risk-Averse Thermal-Aware 2D Macro Placement under Workload Uncertainty

Arsen Kozhabek¹ and Johannes Milz²

1. School of Electrical and Computer Engineering, 2. H. Milton Stewart School of Industrial and Systems Engineering

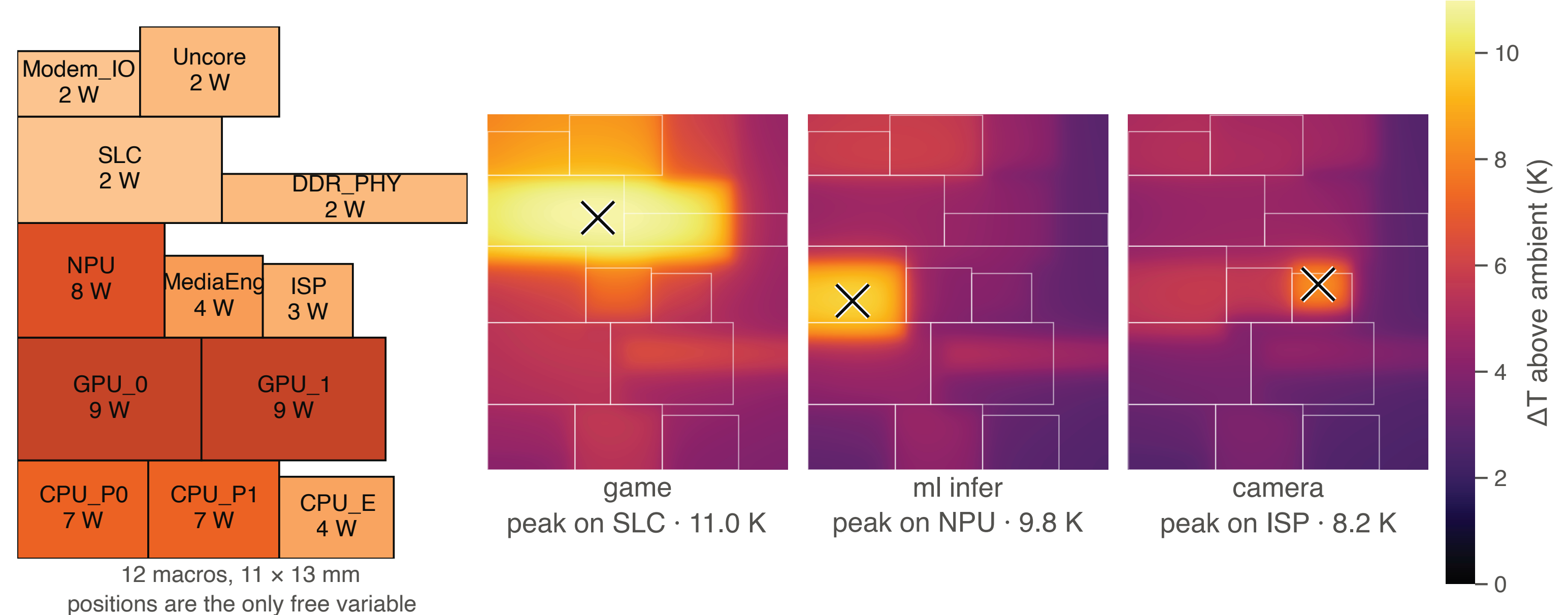


The problem

The workload is uncertain. The placement is fixed.

Placing the macros (a pre-designed functional block of circuitry) of a computer chip is a trade-off between wirelength, timing, and heat.

The thermal component of macros depends on software: different applications and execution phases activate different functional units, producing different power maps across the die, and therefore different hotspots.; The macro placement is chosen once, during the chip-design stage, and cannot change after fabrication. The workload, however, keeps changing, so a single, irreversible placement is judged against an entire ensemble of workloads it was never shown.



Uncertain workloads create a great challenge

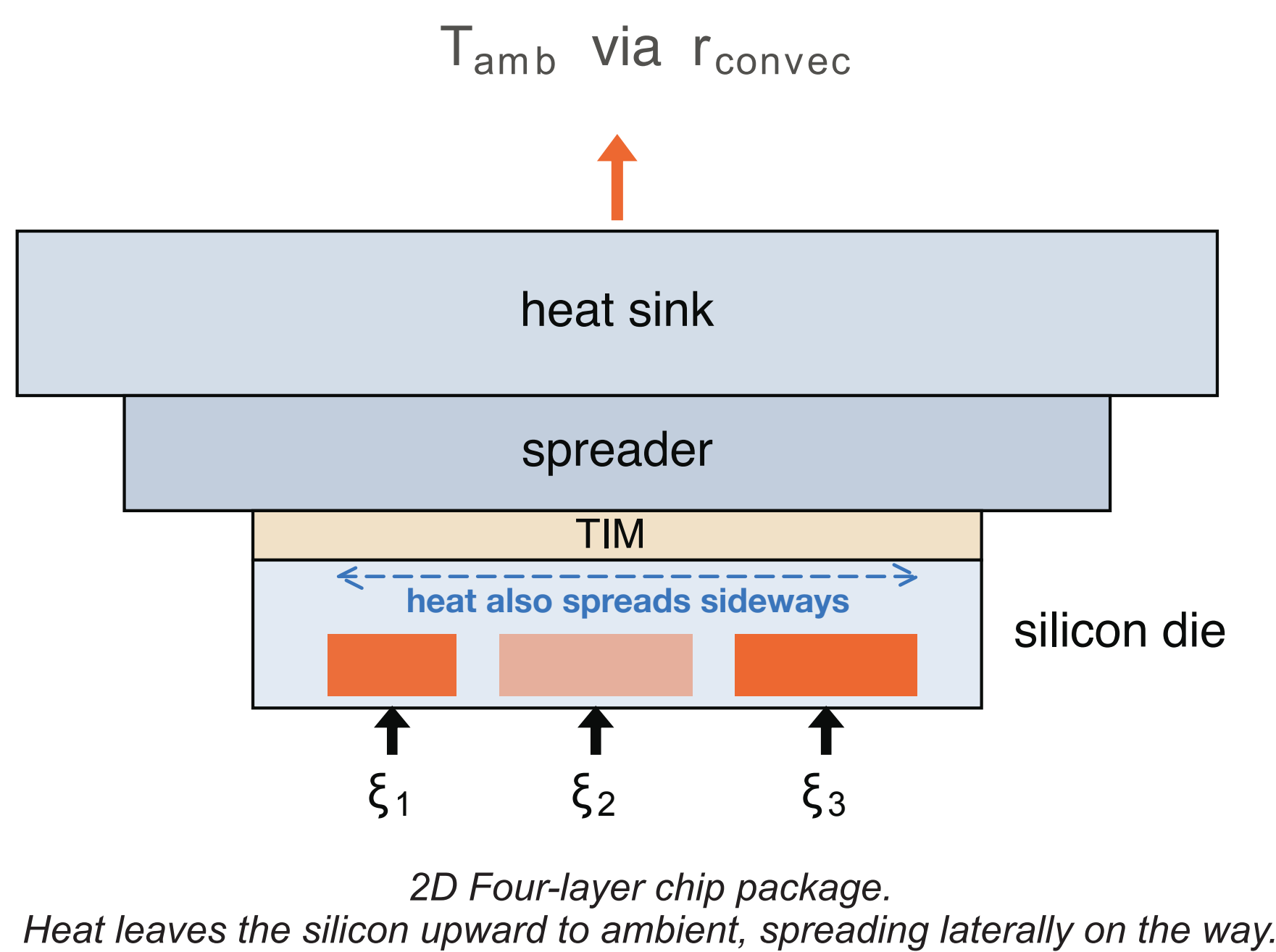
The chip carries **B** rectangular macros of fixed size. Position is the only choice, so a placement is the vector **p** of macro centers:

$$\mathbf{p} = ((c_x^1, c_y^1), \dots, (c_x^B, c_y^B)) \in \mathcal{X} \subset \mathbb{R}^{2B}$$

where **X** is the die rectangle intersected with pairwise non-overlap. A placement alone has no temperature; that requires a workload **ξ**, the activity level of each macro.

Ranking placements needs one scalar, and the binding one is peak temperature. A placement meets a whole ensemble of workloads, so it faces a distribution of peaks; ranking means choosing which statistic of that distribution to minimize.

From placement to temperature



Discretizing the package on a uniform grid turns the physics into one sparse linear system:

$$\mathbf{G} T(\mathbf{p}, \xi) = \mathbf{Q}(\mathbf{p}, \xi)$$

with **G** from geometry and materials by Kirchhoff's current law. Because the layers are laterally uniform, **G** does not depend on where the macros sit — placement changes where heat is injected, not how it propagates, and enters only the right-hand side:

$$\mathbf{Q}(\mathbf{p}, \xi) = \mathbf{Q}_0 + \mathbf{A}(\mathbf{p}) \xi$$

The hottest silicon cell of **T = G⁻¹Q** is the scalar we asked for:

$$\mathbf{M}(\mathbf{p}, \xi) = \max_{i \in I_{Si}} T_i(\mathbf{p}, \xi) - T_{amb}$$

Fixed **G**, placement-dependent **Q** is what makes this possible. One factorization serves every scenario; the derivative of **M** is one back-substitution against it. Thousands of workloads per step become affordable — so an objective over the whole distribution is finally in reach.

The machine and measurement

Mean, or tail?

The mean smooths over precisely the workloads a thermal designer worries about. So alternative (or addition) to it is **CVaR_α**, the mean of the worst **(1-α)** fraction of outcomes:

$$\min_{\mathbf{p}} \mathbb{E}_{\mathbf{p}}[\mathbf{M}] \quad \text{versus} \quad \min_{\mathbf{p} \in \mathcal{X}} \text{CVaR}_{\alpha}[\mathbf{M}(\mathbf{p}, \xi)]$$

$$\text{CVaR}_{\alpha}[\mathbf{M}] = \min_{t \in \mathbb{R}} \left\{ t + \frac{1}{1-\alpha} \mathbb{E}[(\mathbf{M} - t)_+] \right\}$$

At the inner minimizer, **t* = VaR_α** and the t-term drops out, leaving the mean's per-scenario gradients reweighted onto the tail:

$$\frac{1}{1-\alpha} \mathbb{E}[\mathbf{1}\{\mathbf{M} \geq t^*\} \nabla_{\mathbf{p}} \mathbf{M}] \in \partial_{\mathbf{p}} \text{CVaR}_{\alpha}[\mathbf{M}(\mathbf{p}, \xi)]$$

And the tail costs information. That subgradient is supported on only **[N(1-α)]** scenarios at **α = 0.9**, ten times fewer samples than the mean gradient. Noisier, and sharper.

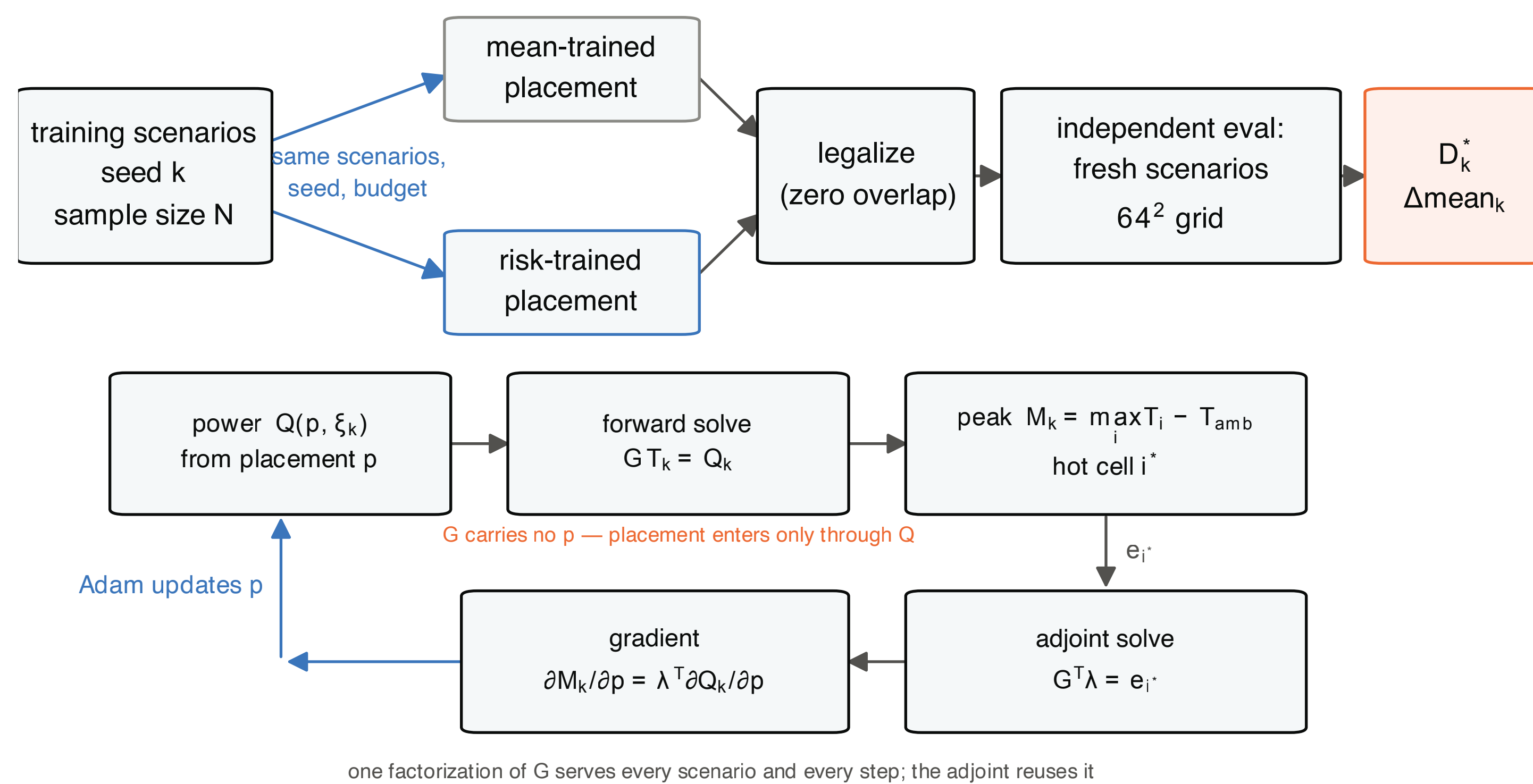
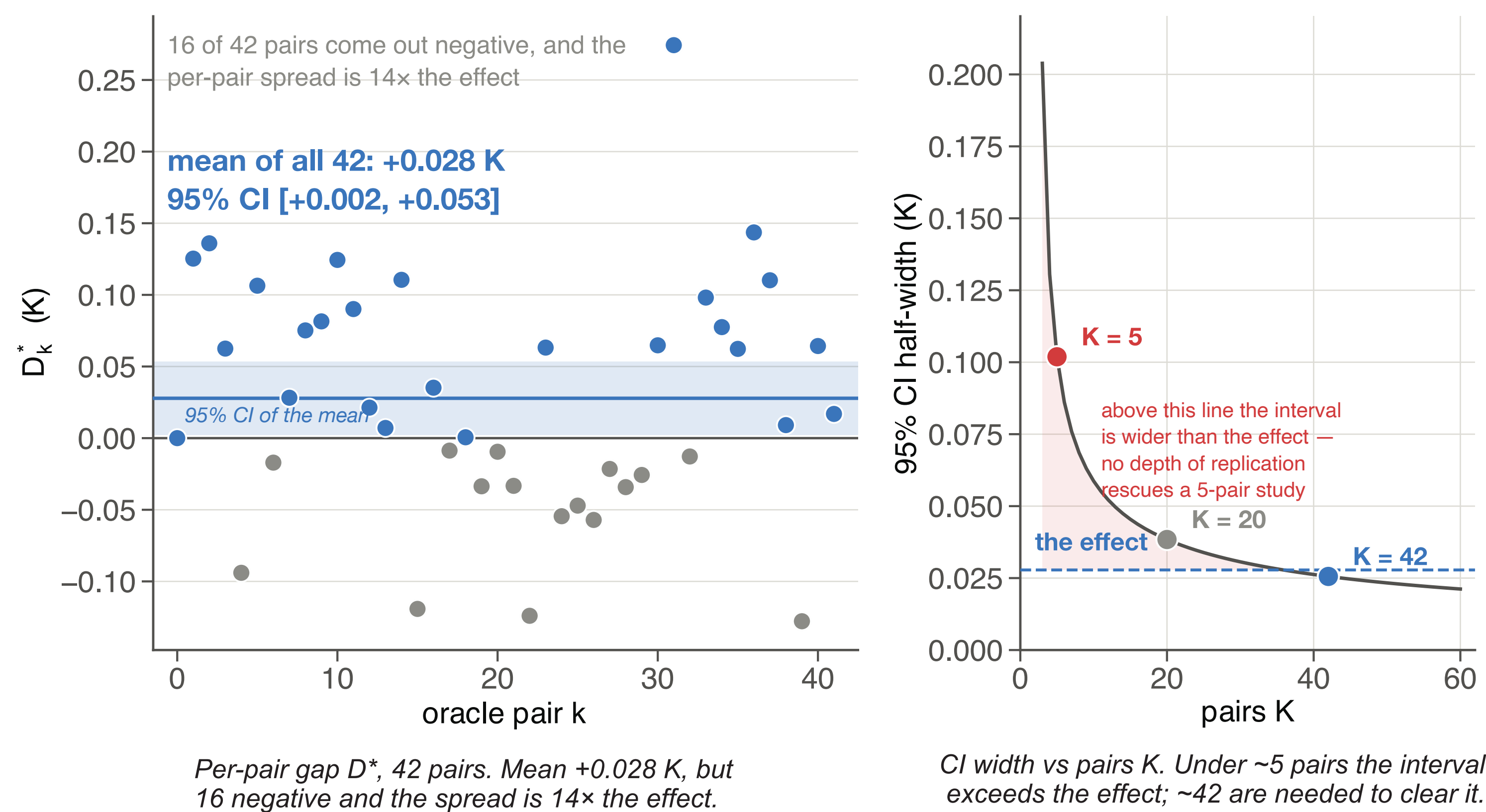
Which makes this a measurement problem

Whether the tail objective actually lowers tail temperature cannot be settled by analysis, only measured. The comparison is not between two placements but between two procedures: the training run that minimizes CVaR against the one that minimizes the mean. A procedure is a random map **ω** ↦ **p(ω)** from seed to placement, so the quantity of interest is a property of that map — the expected gap:

$$\mathbf{D}_k^* = \text{CVaR}_{\alpha}(\mathbf{p}_{\text{mean}}) - \text{CVaR}_{\alpha}(\mathbf{p}_{\text{cvar}}), \quad \bar{\mathbf{D}}^* = \frac{1}{K} \sum_{k=1}^K \mathbf{D}_k^*$$

Each oracle pair **k** gives one measured gap **D_k**; we report their mean **D** over **K = 42 pairs**. Neither arm returns an argmin — **X** is disjunctive, **M(·, ξ)** is nonconvex, the descent is a subgradient method — and each minimizes an empirical **CVaR** on finitely many scenarios. So **D*** compares procedures, not optima, and every property of the

With exact minimizers the true gap is **≥ 0** by construction — the **CVaR-optimal** placement cannot be beaten on **CVaR**. That does not transfer to **D***, so a significantly negative **D*** is diagnostic — of finite samples, or one arm stuck in a worse local optimum, never of the theory. Our first study measured exactly that.

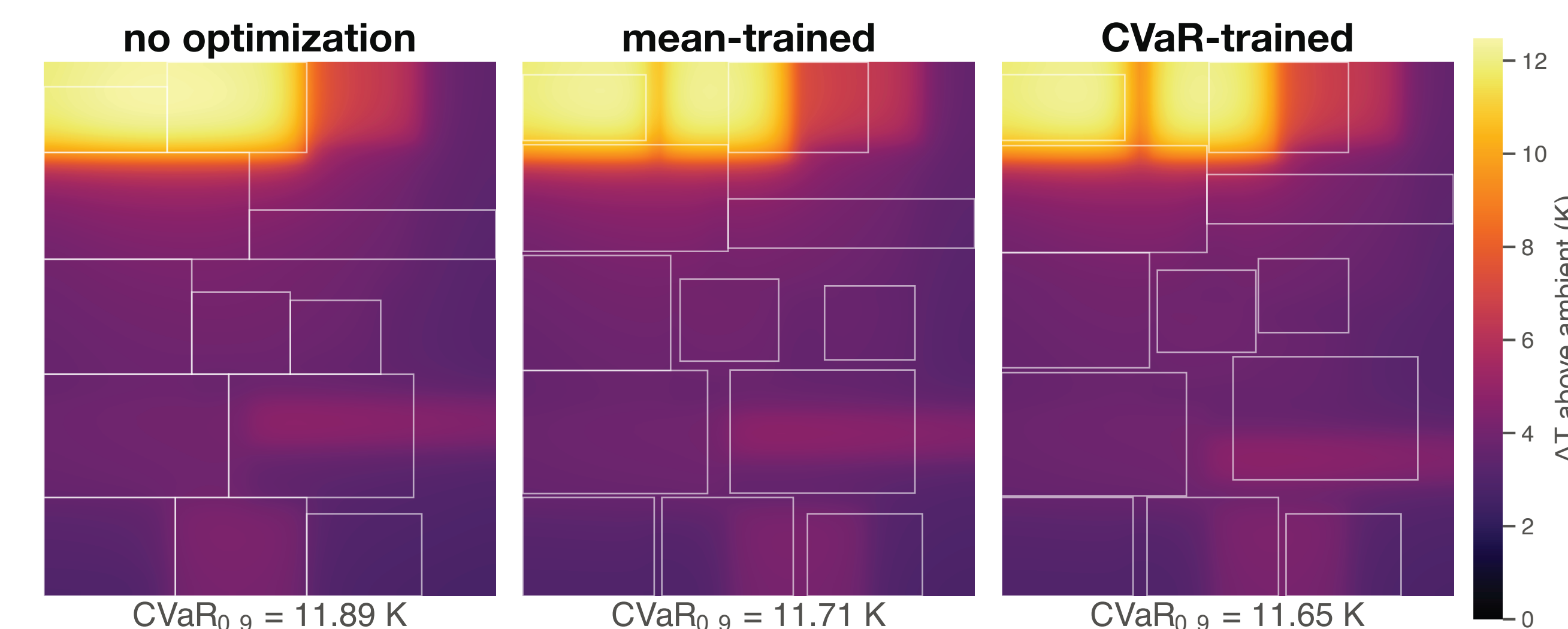


What survives measurement?

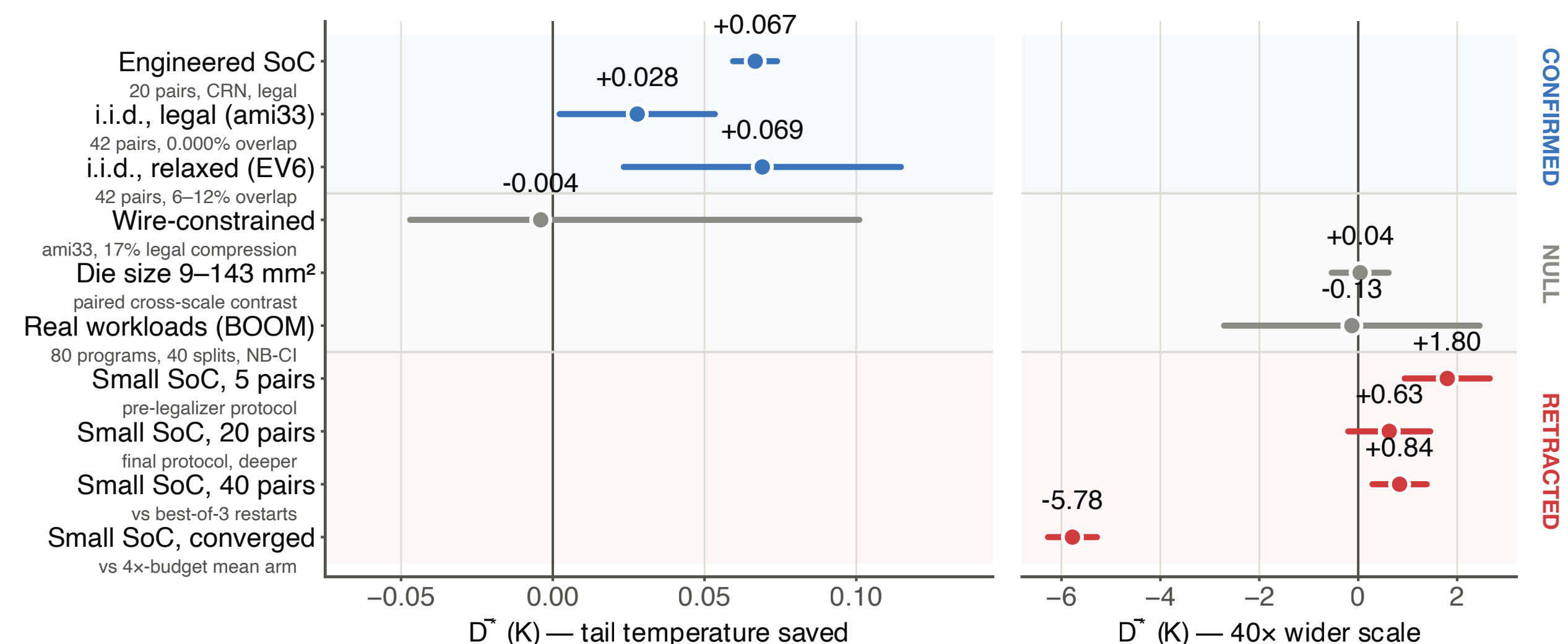
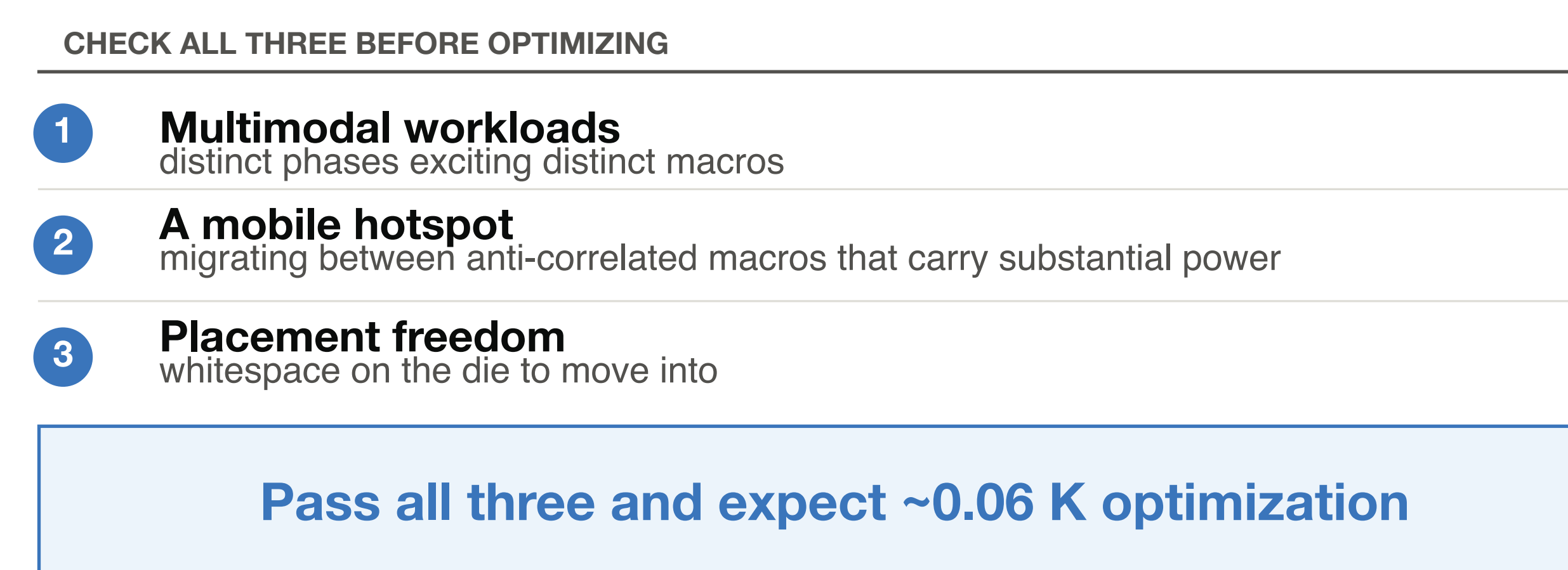
The result

A trained placement is not a measurement — it is an optimizer's output: finitely many Adam steps, a finite sample, a chosen mesh, constraints enforced only approximately. Each is a knob that can favor one arm, and the sharper tail objective exploits any slack. So the knobs are equalized before anything is measured — common random numbers, matched budget, identical legalization — and the mesh is escaped rather than shared: jittered in training, then scored on a finer one neither arm saw.

+0.066 K [+0.059, +0.074] engineered SoC, trade-shaped (Δmean < 0)
+0.028 K [+0.002, +0.053] guaranteed-legal placements under i.i.d. power



When should a chip designer bother?



D* across testbeds and protocols, grouped by verdict. Left: sound protocols — the effect is real but under 0.1 K. Right (40x wider scale): null results, and small-sample or illegal-overlap runs that inflate **D*** and were retracted, not reported.

Acknowledgements

Supported by ISyE 2026 Summer Scholars Program at Georgia Tech. This research was supported in part through research cyberinfrastructure resources and services provided by the Partnership for an Advanced Computing Environment (PACE) at the Georgia Institute of Technology.

Workload traces from the McPAT-Calib release; the ami33 floorplan from the GSRC benchmark distribution. Neither is redistributed with our code.

- References:
- Rockafellar RT, Uryasev S (2000) Optimization of conditional value-at-risk. *J Risk*, 2(3):21-41.
 - Huang W, Ghosh S, Velusamy S, Sankaranarayanan K, Skadron K, Stan MR (2006) HotSpot: a compact thermal modeling methodology for early-stage VLSI design. *IEEE Trans VLSI Syst*, 14(5):501-513.
 - Kingma DP, Ba J (2015) Adam: a method for stochastic optimization. *ICLR*.
 - Lin Y, Dhar S, Li W, Ren H, Khalilany B, Pan DZ (2019) DREAMPlace: deep-learning-toolkit-enabled GPU acceleration for VLSI placement. *DAC*.
 - Zhai J, et al. (2022) McPAT-Calib: a RISC-V BOOM microarchitecture power modeling framework calibrated against RTL sign-off power. *IEEE TCAD*.
 - Di Mauro A, et al. (2022) Kraken: a direct eventframe-based multi-sensor fusion SoC for ultra-efficient visual processing in nano-UAVs. *arXiv:2209.01065*.
 - Holm S (1979) A simple sequentially rejective multiple test procedure. *Scand J Statist*, 6(2):65-70.
 - Nadeau C, Bengio Y (2003) Inference for the generalization error. *Machine Learning*, 52(3):239-281.
 - MCNC/GSRC floorplanning benchmark suite (ami33); 33 macros, 78 nets.
 - DIFFChip: differentiable thermal-aware chip placement. *arXiv:2502.16633* (2025).
 - ATPlace2.5D: analytical thermal-aware placement for 2.5D systems. *ICCAD* (2024).