

基于 60+15+15 分阶段学习的汽车热管理云标定优化

指标残差校正、初态分区与配对反馈验收

摘要

汽车热管理标定同时涉及温度跟踪、恢复过程、区域平衡和执行器代价。有限运行记录难以支撑高维参数的直接寻优，预测误差也可能在候选选择中被放大。本文研究一种基于既有可行策略库的分阶段适配方法：以 60 个互斥模拟工况建立指标响应代理，以 15 个工况学习低维残差，再以 15 个工况的配对模拟输出实施多指标验收和基线回退。优化目标由逐工况归一化指标的均值与 90% 分位数组成，标定参数满足量化、边界及锁定约束，完整分区策略仅在试验初始时读取允许特征。六个上下文的新适配试验共执行 4410 次候选工况评估，四种方法在冻结后接受 720 个未参与开发的测试工况检验。完整反馈在 4 个上下文通过全部测试保护规则，其中 2 个保持基线，另 2 个兼具正收益，改善为 1.2360% 和 3.6809%。其余 2 个上下文的综合目标虽下降，稳态波动保护未通过；残差校正未改善六个测试集的指标预测误差。另一个独立的分区策略确认试验在 12 个面板上取得 0.3650% 至 1.2659% 的综合目标改善。两组试验共同用于分析有限策略库适配与反馈验收的作用。研究结论限于登记模拟模型与工况，执行器努力度为代价代理，不代表实车节能率。

关键词：汽车热管理；云标定；分阶段学习；残差校正；受约束优化；配对反馈

代码与数据 (Code and data)：实验材料尚未公开发布；数据范围与可用性见文末声明。

1 引言

电动汽车的热管理系统需要协调座舱温度、电池热状态及部件运行需求。控制性能随环境、初始状态和设定温度计划改变，同一参数调整可能降低平均误差，同时增加某些转场的等待时间或局部波动。集成模型预测控制研究已将多个热管理子系统的温度要求与功耗纳入联合决策 (Wang et al., 2024)；面向舒适性的控制研究也表明，温度跟踪指标与资源代价之间需要明确的权衡 (Kwak et al., 2023)。这些工作说明了多目标建模的重要性，也提示标定结果应接受跨工况和约束层面的检验。

云标定所面对的实际问题往往出现在完整物理模型形成之前：数据包含匿名测点、参数快照和少量运行记录，参数可调范围受工程约束限制，模型与执行结果之间仍有偏差。单次参数快照对应的观测轨迹，不能唯一识别全部参数的干预响应。因此，本文将任务限定为利用已有可行策略库完成新工况适配。策略库提供完整参数向量和初态分区规则，新增数据用于估计各候选的指标表现、校正预测偏差并检验候选的实际模拟响应。

本文以“60+15+15”组织适配过程。60 个工况用于基础指标建模；随后 15 个工况用于独立的残差学习；最后 15 个工况用于反馈验收及最终策略选择。三个集合互斥，每个工况可以对多个策略进行配对评估。这里的数量表示模拟工况数，单个工况持续 120 min；它们不表示 90 天实车采集记录，也不包含既有策略库的形成成本。独立测试只在三个阶段完成并冻结策略后执行。

本文的工作包括三个方面。首先，给出从参数约束、初态路由、指标响应拟合到反馈验收的完整数学定义，使模型输出与可执行标定策略之间的关系明确。其次，以共享候选库比较基础代理、残差代理和完整反馈三种选择方式，分别检查阶段外预测质量与控制结果，避免将训练拟合度直接解释为优化能力。最后，将新适配试验与另一组冻结分区策略确认分开报告，分析小样本反馈、平均性能和约束保护之间的边界。本文不讨论实时在线控制器的替代，也不将大模型解释直接作为数值验收依据。

2 相关研究

2.1 控制导向模型与热管理优化

热管理优化通常需要将动态模型、状态约束和控制代价结合。Wang et al. (2024) 研究集成热管理的模型预测控制，Kwak et al. (2023) 将热舒适需求纳入车辆空调控制。快速充电场景下，功率与热管理的联合约束又引入不同的预测和鲁棒性要求 (Hu et al., 2023)。本文沿用“多指标且有明确边界”的建模原则，但研究对象是离线标定参数与有限策略库的适配，不采用上述工作的在线滚动控制问题，也不将其能耗结果用于本研究的效果比较。

2.2 代理模型、残差校正与受约束搜索

岭回归通过二次正则稳定相关输入下的系数估计 (Hoerl and Kennard, 1970)，适合构造维度受控的指标代理。模拟器校准研究区分参数校准与模型偏差 (Kennedy and O'Hagan, 2001)，为基础模型外另设残差层提供了方法动机。本文使用确定性的低维岭残差，不采用贝叶斯校准中的高斯过程不确定性模型，也不继承其推断性质。

一般黑箱优化可通过函数值、代理模型和约束检查组织搜索 (Conn et al., 2009)。受约束直接搜索的理论还需要网格、搜索方向及可行性等条件 (Audet and Dennis, 2006)。本文在固定有限候选库上进行枚举评分和验收，不将这一实现称为网格自适应直接搜索或完整信赖域算法。其可解释性来自有限集合、明确评分及可审查的回退规则；关于连续参数空间或全局最优的结论不在本文范围内。

2.3 大模型与数值评价器协同

大语言模型可以读取候选及其评价，提出下一批优化建议 (Yang et al., 2024)，也可参与贝叶斯优化中的初始化、代理构造或候选采样 (Liu et al., 2024)。这些研究支持将语言模型视为外层建议模块，具体数值和约束仍需要任务评价器。本文已有策略库的选择环节包含真实保存的大模型调用；新增三阶段适配试验完全使用固定计算程序，不额外调用模型。因此，实验比较用于分析指标代理与反馈验收，不构成大模型相对于同预算数值优化器的增益证明。

3 问题定义与优化方法

3.1 可行参数、完整策略与合法输入

设上下文 $c \in \{1, \dots, 6\}$ 表示一组固定的模拟参数敏感度配置， $p^0 \in \mathbb{R}^{88}$ 为参考标定快照。选定 16 个可调整分量，记归一化增量 $z \in \mathbb{R}^{16}$ 。每个分量通过尺度 d_j 映射并量化为

$$p_j(z_j) = Q_j(p_j^0 + d_j z_j), \quad \mathcal{Z} = \{z \in [-1, 1]^{16} : p_j^- \leq p_j(z_j) \leq p_j^+, |p_j(z_j) - p_j^0| \leq d_j\}. \quad (1)$$

Q_j 按预先给定的原参数步长四舍五入并裁剪到边界， d_j 由允许改变量确定。本次选定参数的双侧改变量尺度相同，其余 72 个分量保持 p_j^0 。有效候选必须以完整向量满足约束，不能分别挑选不同方案的单项“最好值”组成未经评估的新向量。

试验初始输入为

$$x = (e_0, |\Delta T_0|, \Delta r_1, V_r), \quad V_r = \sum_t |r_{t+1} - r_t|. \quad (2)$$

其中 e_0 为初始均温相对目标的偏差， $|\Delta T_0|$ 为初始双区温差幅值， Δr_1 和 V_r 来自试验前已给定的目标计划。模型族、样本编号、随机种子、隐藏扰动系数及试验后的响应指标均不属于路由输入。一个完整策略 π_a 将 x 映射到 $\mathcal{Z}$ ，且该向量在整个工况内保持不变。

候选既包括常数向量，也包括初态分区规则。以恒定目标计划下的三分区为例，取一个允许特征 x_j 和预先确定的阈值 $a < b$ ，定义

$$\pi(x) = \begin{cases} 0, & V_r > 0, \\ z_L, & V_r = 0, x_j \leq a, \\ z_M, & V_r = 0, a < x_j \leq b, \\ z_R, & V_r = 0, x_j > b. \end{cases} \quad (3)$$

各叶子都是经量化验证的完整 16 维向量。二分区使用一个阈值，其他候选可采用同一组合法特征的条件合取。零向量策略 π_0 始终保留。新增适配试验使用固定候选集合 $\mathcal{P}_c = \{\pi_0, \dots, \pi_{K_c-1}\}$ ，不在测试结果出现后修改阈值或叶子。

3.2 配对指标与均值-分位数目标

对于同一工况 i ，候选 a 与基线使用相同设定计划、初态和噪声种子，得到指标 m_{iak} 与 $b_{ik} = m_{i0k}$ 。表1给出 8 个加权指标及归一化下限。温度跟踪误差为完整 121 时点上均温与目标的绝对误差均值；稳态误差和波动在末 1200s 窗口上计算，包含两个端点共 21 点，标准差采用总体形式。执行器努力度取模拟执行量平方除以效率代理，并报告其时均值与峰值。

表 1: 综合目标中的指标、固定权重与同工况归一化下限。另有热舒适及区域温差违反率作为保护指标。

指标	权重 w_k	下限 f_k	单位
跟踪绝对误差均值	0.18	0.10	°C
稳态绝对误差均值	0.14	0.10	°C
稳态跟踪误差标准差	0.08	0.10	°C
进入并保持目标带的等待时间	0.10	60.00	s
分段截尾等待比例均值	0.18	0.05	无量纲
可评估目标段恢复失败率	0.20	0.10	比例
执行器努力度均值	0.08	0.05	代理值
执行器努力度峰值	0.04	0.05	代理值

恢复要求在同一恒定目标段内，均温跟踪误差处于 $\pm 1^\circ\text{C}$ 内，且连续采样跨度达到 300s；在 60s 采样间隔下需要连续 6 个点。整段等待指标取首次满足条件的连续段起始时刻，未出现时截尾到观察终点。分段恢复按各目标平台段分别计算；每个可评估段同时保留是否达标、截尾等待比例及失败标志。过短的不可评估段单独记录，温度误差与舒适性仍参与评价。热舒适及区域温差违反分别采用绝对跟踪误差严格大于 2°C 、双区温差幅值严格大于 2°C 的判据。保护规则还要求基线已成功的目标段在候选下继续成功，防止平均温度误差改善掩盖个别转场的恢复丢失。

令同工况归一化指标、综合分数和集合目标分别为

$$r_{iak} = \frac{m_{iak}}{\max(|b_{ik}|, f_k)}, \quad s_{ia} = \sum_{k=1}^8 w_k r_{iak}, \quad (4)$$

$$J_D(a) = 0.7 \text{mean}_{i \in D}(s_{ia}) + 0.3 Q_{0.9}(s_{ia}), \quad G_D(a) = 100 \frac{J_D(0) - J_D(a)}{J_D(0)}. \quad (5)$$

经验分位数使用线性插值。 $G_D > 0$ 表示该固定综合目标下降。均值与高分位共同反映典型工况和较大损失工况，所用 $Q_{0.9}$ 是分位点；它与尾部条件期望型风险度量具有不同定义 (Rockafellar and Uryasev, 2000)，本文不赋予该混合目标 CVaR 的性质。

所有候选使用同一组基线分母。先计算完整候选的分数分布，再比较 $J_D(a) - J_D(0)$ ；一般而言， $Q_{0.9}(s_a - s_0)$ 并不等于 $Q_{0.9}(s_a) - Q_{0.9}(s_0)$ 。同样，分区策略必须先逐工况选择完整参数叶，再计算整体指标分布，不能通过各叶子摘要的线性加权替代尾部分位数。

3.3 60 工况基础拟合与 15 工况残差学习

记基础集、残差集与反馈集为 D_0, D_1, D_2 ，其工况数依次为 60、15、15，且两两不交。所有候选在每个工况上进行配对评估，形成 $K_c \times |D_j|$ 个策略响应。候选共享工况，因此重复策略探测不增加独立工况数。

以 D_0 的四特征均值与标准差进行标准化，得到 $q = (x - \mu)/\sigma$ ；接近零的标准差下限为 10^{-8} 。基础特征为 $\phi(x) = (1, q_1, q_2, q_3, q_4)^\top$ 。将所有候选的 8 项指标比值并列为响应矩阵 $Y_0 \in \mathbb{R}^{60 \times 8K_c}$ ，基础模型通过

$$B^* = \arg \min_B \|Y_0 - \Phi_0 B\|_F^2 + \lambda \|LB\|_F^2, \quad L = \text{diag}(0, 1, 1, 1, 1) \quad (6)$$

估计。截距不受惩罚。候选正则系数为 $\{0.01, 0.1, 1, 10\}$ ，在 60 工况内部按“前 30 预测后 10、前 40 预测后 10、前 50 预测后 10”三个扩展块选择。四特征的预处理统计量在基础阶段固定，后续两个阶段及测试集均不重新估计。这里拟合的是完整工况的指标响应，不将其解释为逐时刻温度动态模型。

基础模型在读取 D_1 的模拟响应前保存。在 D_1 上，残差目标为 $E_1 = Y_1 - \Phi_1 B^*$ ，采用低维特征

$$\psi(x) = (1, q_1^2, q_2^2, \mathbf{1}[V_r = 0])^\top, \quad C^* = \arg \min_C \|E_1 - \Psi_1 C\|_F^2 + \|L_1 C\|_F^2, \quad (7)$$

其中 $L_1 = \text{diag}(0, 1, 1, 1)$ ，残差正则固定为 1。该设计将偏差校正限制为初始温差曲率、区域温差曲率及恒定计划效应，控制 15 个工况下的模型维度。最终预测为

$$\hat{r}_a(x) = [\phi(x)^\top B_a^* + \psi(x)^\top C_a^*]_+, \quad (8)$$

其中 $[\cdot]_+$ 逐元素截断到非负值。基础代理消融去掉残差项。非负投影只是数值范围处理，不保证对候选性能的保守预测。

阶段外预测评价采用

$$E_D = \frac{1}{K_c} \sum_{a=0}^{K_c-1} \sum_{k=1}^8 w_k \sqrt{\frac{1}{|D|} \sum_{i \in D} (\hat{r}_{iak} - r_{iak})^2}. \quad (9)$$

该值是归一化指标比值的加权 RMSE，无温度单位。基础模型在 D_1 上评价后才学习残差；基础与校正模型在 D_2 上的评价都发生于任何反馈更新之前。不能用 D_0 或 D_1 训练集上的拟合误差代替上述预测评价。

3.4 15 工况配对反馈、验收与冻结

两个代理消融仅根据 D_2 的初态和预设计划计算预测目标 $\hat{J}_{D_2}(a)$ ，在读取候选的反馈响应前分别冻结 $\arg \min_a \hat{J}_{D_2}(a)$ 。完整方法随后使用 15 个工况的直接模拟结果，建立可接受集合

$$\mathcal{A}_{D_2} = \{a : H_{D_2}(a) = 1, J_{D_2}(a) \leq J_{D_2}(0)\}, \quad a^* = \arg \min_{a \in \mathcal{A}_{D_2}} J_{D_2}(a). \quad (10)$$

相同目标值按固定候选顺序处理。基线自身始终可接受，因而集合非空。 H_D 同时检查总体及两个模拟模型族的保护条件：10 项指标均值不增加，数值容差 10^{-9} ；各指标 90% 分位数增加不超过 1%；基线已成功目标段不丢失；工况及目标段的截尾计数不增加；每个工况至少包含一个可评估目标段。均值约束与分位数约束分别判断，不用综合目标补偿某个保护指标的退化。

主实验：每个上下文使用互斥模拟工况；每工况为 120 min

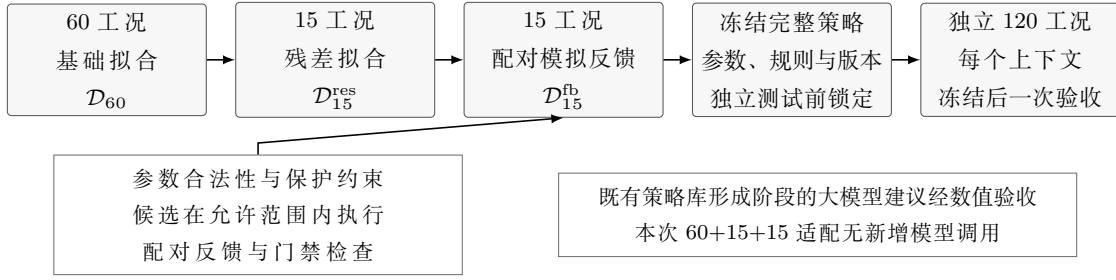


图 1：以互斥工况组织的 60+15+15 适配流程。每个工况可以评估多个完整策略。独立测试在模型与各方法选择均冻结后开放；其输出不返回开发阶段。

表 2：有限策略库上的分阶段适配算法。

输入：固定候选库 $\mathcal{P}_c$ 、参数约束、互斥 D_0, D_1, D_2 及未开放的 D_H 。

输出：一个冻结的完整策略及独立测试结果。

1. 校验所有参数叶及路由特征，保留完整零向量基线。
2. 在 D_0 配对评估全部候选；构造指标比值，选择正则并拟合 B^* 。
3. 在 D_1 先生成基础预测，再读取响应并拟合 C^* 。
4. 在 D_2 只读取初态与计划，冻结基础代理及残差代理的候选选择。
5. 配对评估 D_2 的候选，按式(5)评分，逐项执行 H_{D_2} 。
6. 按式(10)选择反馈候选；无可接受改进时保留基线。
7. 冻结所有上下文的模型、候选选择及参数；一次开放 D_H 并统一评价四种方法。
8. 报告目标、各保护指标、启用覆盖与预测误差；独立测试不触发再次选择。

有限集合性质与计算预算。 在固定 D_2 、固定指标和候选库下，式(10)给出可接受集合内的最小观察目标，且 $J_{D_2}(a^*) \leq J_{D_2}(0)$ 。其依据是基线可行和有限集合枚举；这不推出未见工况上的逐项不退化，也不推出连续参数空间的最优性。一次完整研究的模拟调用量为 $\sum_c K_c(60 + 15 + 15 + N_H)$ 。矩阵拟合只有 5 维基础特征和 4 维残差特征，主要开销来自模拟函数评价。本研究为保持对照完整而计算全部候选，代理排序没有减少该枚举预算；扩大策略库后的采样效率需另行研究。

3.5 独立测试及大模型接口

独立集 D_H 在所有上下文的四方法选择冻结后生成。不同策略共享同一工况和噪声实现，重采样也以工况对为单位。对候选与基线分数同时抽取相同索引，分别重算两侧目标并形成差值分布，采用 10000 次 bootstrap 给出近似区间 (Efron, 1979)。主适配试验报告每上下文的探索性 95% 双侧差值区间，不据此声明多个上下文的联合显著性。

若开展重复确认，检验预算应在查看新结果之前安排 (Lan and DeMets, 1983)。本文辅助确认使用其预先固定的单侧分位点 0.996875，同时核对常规单侧 0.95 分位点。预算分配与高分位重采样均不能自动建立有限样本覆盖保证，本文只按明确的登记规则描述计算结果。

大模型接口的允许输出是已完成开发资格检查的候选标识及选择解释，输出标识必须属于当前候选集合。数值计算模块验证标识、完整参数、策略规则和复算结果。模型既不读取独立测试后再修改选择，

也不通过自然语言指令覆盖参数边界。已有候选比较采用 DeepSeek 的高思考配置，所选策略随后完成 7200 次开发配对模拟复验。本轮适配使用冻结库的确定性计算，不产生新的大模型请求。

4 实验设计与结果

4.1 数据来源、模拟模型与两组实验身份

观测锚点来自一次脱敏运行：151 个时点，间隔 2s，总时长 300s，配套 88 项参数快照。该记录用于约束温度尺度和噪声一阶差分尺度，参数边界来自配套标定约束。两种模拟族的噪声自回归系数分别固定为 0.72 和 0.82，属于模拟假设。模拟敏感度同样是显式建模假设，未由这一单次记录识别为参数因果效应。每个计算工况持续 120 min，以 60s 步长产生 121 个时点。

模拟域包含双区一阶热响应与双热容热响应。前者通过比例-积分控制、执行延迟和速率限制驱动双区冷却；后者加入壁面热状态、温度相关制冷能力、迟滞、斜坡目标和重尾自相关扰动。统一指标从输出轨迹重新计算，避免两种模型内部摘要的窗口或到温定义差异。六个上下文是六组敏感度配置，不表示六辆实车。模型族仅用于工况配额和分族验收，不输入策略路由或指标代理。两种模型族都在适配阶段出现，独立测试检验登记混合分布中的新工况，不属于未见模型族测试。

主实验每上下文包含 60 个基础工况（两族各 30）、15 个残差工况（8 与 7）、15 个反馈工况（7 与 8），冻结后另取 120 个测试工况（两族各 60）。同一上下文内，基础、残差、反馈与测试的随机种子互不相交。不同上下文之间有 5 对整数噪声种子复用，对应的模型族、完整工况及敏感度矩阵均不同，没有完整工况重复；因此只计算上下文内的配对区间，不将跨上下文结果合并为独立同分布样本。候选库由基线、一个既有分区策略及其不同非零完整叶的常数策略组成，去重后的 K_c 依次为 5、3、3、4、3、3。六个上下文累计 540 个开发工况与 720 个未参与开发的测试工况；开发调用 1890 次，测试调用 2520 次，总计 4410 次。

这些 90 工况用于已有策略库的适配，不包含该库此前的参数搜索、分区筛选和模拟投入。另一组辅助确认采用不同开发库筛选与冻结流程，在六上下文的 A/B 两个面板上各评价 120 个工况，共 1440 个独立确认工况。两组实验的样本和目的分开核算，辅助确认的收益不归属为新 90 工况适配从零得到的结果。

4.2 对照方法与阶段外预测质量

四种方法分别为固定基线、仅基础代理选择、基础加残差代理选择、配对反馈选择。两个代理消融在阅读反馈候选真实响应之前冻结；完整方法使用这些响应执行约束验收。所有方法共享相同库和配对工况，差异在于使用哪些阶段的信息。该设计比较的是信息和验收模块的作用，不是四个互相独立的同成本端到端寻优器。

表3显示，残差校正仅在上下文 1 的反馈集上降低预测误差，在六个测试集上均未降低该误差。因而本次单次划分不支持“15 工况残差学习改善样本外预测”的结论。该现象与小样本校正的不稳定性相容，但当前设计无法将原因唯一归于样本量、特征形式或模型族混合。基础代理的测试误差为 0.1177 至 0.1314，校正后为 0.1201 至 0.1439；这些数值描述指标比值预测，不能换算为温度轨迹拟合优度。

4.3 新适配试验的独立结果

表4中，完整反馈在上下文 1 和 2 保留基线，在其余四个上下文取得正的综合目标改善，但测试保护状态为 4/6 通过。其中上下文 5 和 6 同时具有正收益和全部保护通过，改善分别为 1.2360% 和 3.6809%。上下文 5 的反馈策略比代理选择的常数策略更保守，测试收益也较小；当前结果没有表明反馈选择在每

表 3: 阶段外指标预测误差。数值为式(9)的无量纲加权 RMSE，越小越好。反馈集每上下文 15 个工况，测试集每上下文 120 个工况。

上下文	反馈集，读取响应前预测		冻结后的测试集	
	基础代理	加残差代理	基础代理	加残差代理
1	0.1138	0.1039	0.1177	0.1201
2	0.1119	0.1227	0.1314	0.1397
3	0.0961	0.1267	0.1194	0.1201
4	0.1176	0.1604	0.1216	0.1439
5	0.1274	0.1546	0.1248	0.1324
6	0.0988	0.1262	0.1281	0.1389

表 4: 新适配试验的测试综合目标改善 (%) 与全部保护规则状态。每格同时报告改善和保护状态；基线改善为 0 且在六个上下文均通过。

上下文	仅基础代理	基础加残差	完整反馈
1	2.0523 / 未通过	2.0523 / 未通过	0.0000 / 通过
2	0.2151 / 通过	1.9473 / 未通过	0.0000 / 通过
3	1.0603 / 未通过	1.0603 / 未通过	0.4294 / 未通过
4	4.8573 / 未通过	4.8573 / 未通过	4.7324 / 未通过
5	4.9904 / 通过	4.9904 / 通过	1.2360 / 通过
6	3.6809 / 通过	3.6809 / 通过	3.6809 / 通过
保护通过数	3/6	2/6	4/6

表 5: 完整反馈方法的冻结策略、启用覆盖及探索性 95% 差值区间。 $\Delta J = J(a^*) - J(0)$ 为无量纲目标差；负值有利。区间以每上下文的 120 个工况对重采样，不作跨上下文联合显著性解释。

上下文	策略类型	启用数	改善 (%)	ΔJ 的 95% 区间
1	基线回退	0/120	0.0000	[0.00000, 0.00000]
2	基线回退	0/120	0.0000	[0.00000, 0.00000]
3	初态分区	29/120	0.4294	[-0.00525, -0.00270]
4	常数参数	120/120	4.7324	[-0.05011, -0.03801]
5	初态分区	33/120	1.2360	[-0.01467, -0.00799]
6	常数参数	120/120	3.6809	[-0.04065, -0.02757]

个上下文都优于基础代理。

上下文 3 的完整策略使总体稳态波动均值从 0.1551350 升至 0.1551589 °C，增量为 2.38×10^{-5} °C，超过预设的均值不增加容差；其一阶族对应均值也增加。上下文 4 的双热容族稳态波动 90% 分位数从 0.650973 升至 0.663916 °C，增加 1.9882%，超过 1% 上限。这些数值分别说明严格均值保护对微小变化的敏感性，以及平均改善不能排除尾部波动上升。两者均按原保护规则报告为未通过，不因综合目标下降而改判。

表5的负差值区间支持所选固定策略在相应登记分布下的综合目标改善，但不能替代逐项保护。完整反馈共在 302/720 个测试工况启用非零参数，448 个保持基线；覆盖和收益需要同时解释。基础代理、残差代理与完整反馈分别在 3、2、4 个上下文通过全部测试保护规则；该单次划分的计数只作描述，不建立方法间总体优越性的统计结论。

4.4 冻结分区策略的辅助确认

辅助确认先固定各上下文的完整分区规则及参数叶，再开放两个不同工况面板。资格检查包括开工况总体及分组指标；实际选择后还进行了配对模拟复算。表6报告独立面板结果，图2分别展示改善百分比与候选减基线的绝对目标差区间。

表 6: 辅助确认结果。 ΔJ 为候选减基线的无量纲综合目标差；所列为登记单侧上界。

Context	Panel	Gain (%)	Registered upper ΔJ	Active/total
C1	A	0.3650	-0.001921	35/120
C1	B	0.3898	-0.002015	34/120
C2	A	0.5835	-0.000860	9/120
C2	B	0.5793	-0.000829	10/120
C3	A	0.4887	-0.002472	29/120
C3	B	0.5331	-0.002853	33/120
C4	A	0.9347	-0.005133	35/120
C4	B	0.8136	-0.004307	31/120
C5	A	0.9335	-0.004869	26/120
C5	B	1.2659	-0.007114	34/120
C6	A	0.4072	-0.001690	18/120
C6	B	0.3803	-0.001562	18/120

Auxiliary confirmation: 12 panels

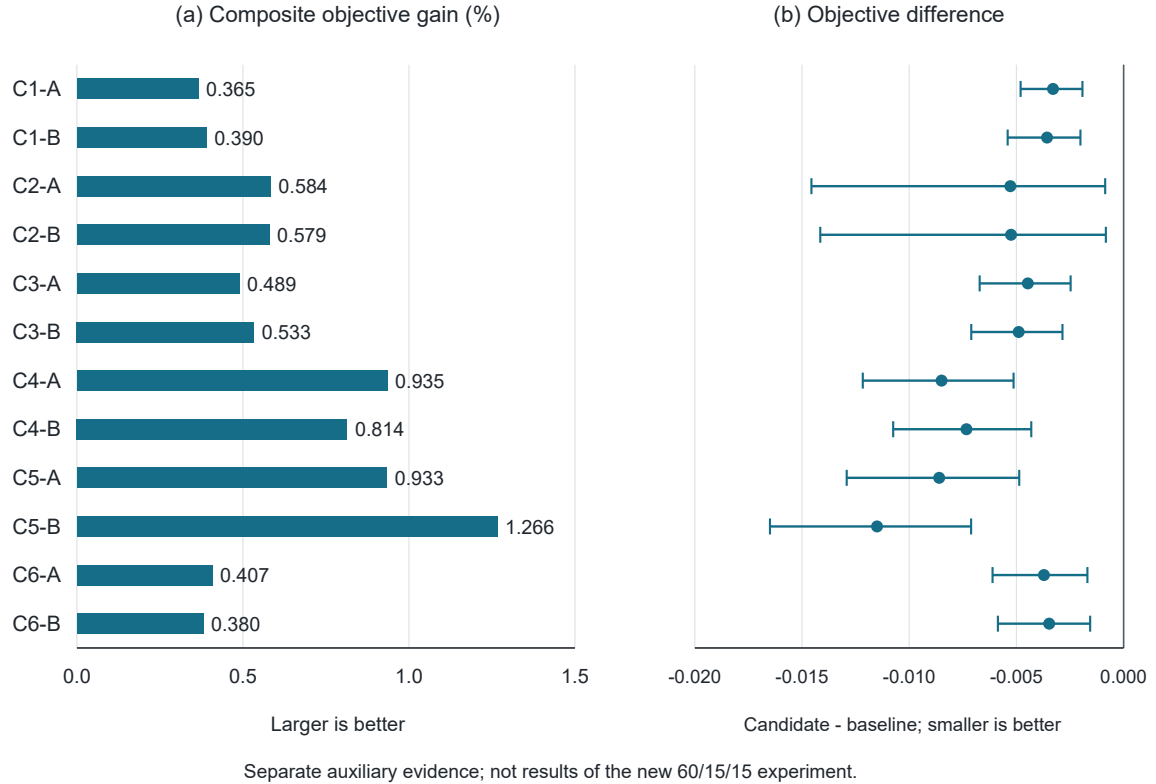


图 2: 辅助确认的 12 个面板。左列为综合目标相对改善，右列为分别计算的登记 bootstrap 单侧上下界；两列的单位不同。每侧采用 99.6875% 分位条件，连线不标为同置信水平的双侧区间。右列负值表示候选目标低于基线。

12 个面板均满足既定 0.15% 最低改善要求和全部保护规则，综合目标改善范围为 0.3650% 至 1.2659%。1440 个工况中 312 个启用非零参数，其余 1128 个保持基线；双热容族 720 个工况全部保持基线，收益为 0。这一覆盖结构意味着改善来自策略实际启用的工况，不能将面板通过解释为每个工况均获得正收益。

C5-B-022: one active condition (median selection)

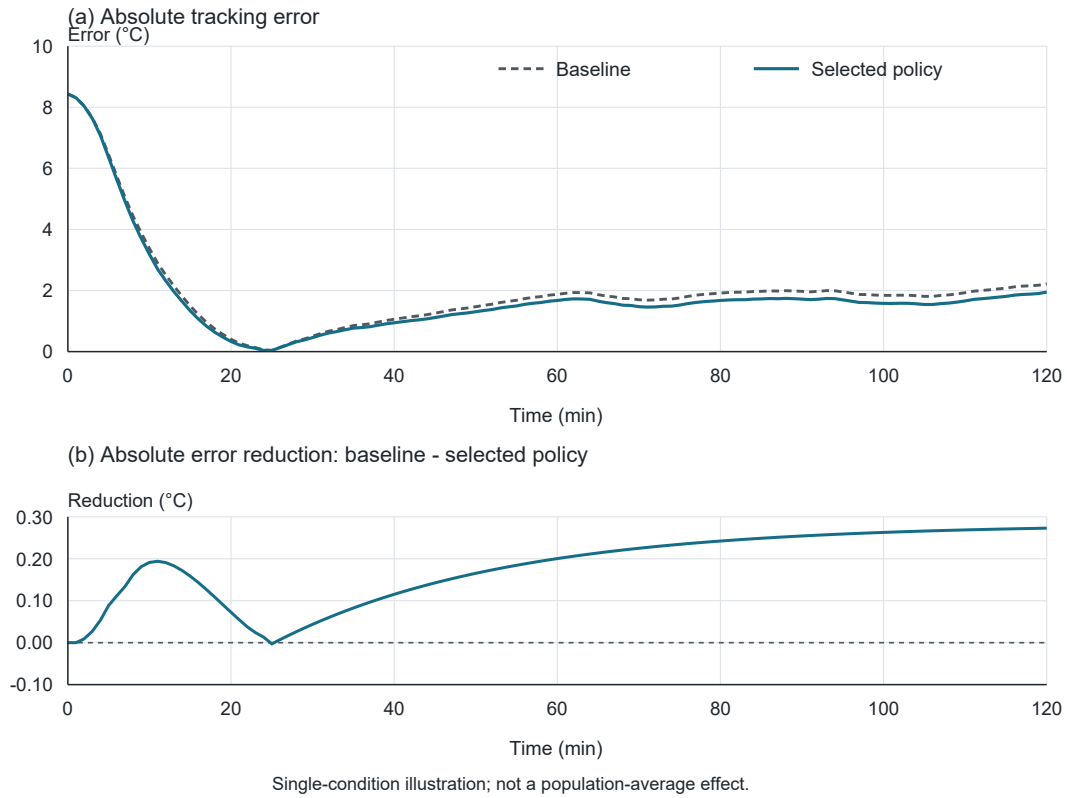


图 3: 辅助确认中上下文 5、面板 B 的一个典型启用工况。示例按本面板启用工况的跟踪 MAE 改善排序取下中位数。上图为共用坐标轴的绝对误差，下图为基线减候选误差；全部 121 个原始采样点保留。该例 MAE 由 1.971409 降至 1.791498 °C，仅代表一个工况。

图3展示改善随时间出现的位置。差值曲线用实际温度误差单位表达微小变化，与综合目标的百分比保持区分。该例不按最大收益筛选，也不承担总体效果估计。主适配试验对保存响应进行了目标、门禁和区间的独立算术核验；辅助确认还按相同工况和参数进行模拟重放。两种检查分别支持计算一致性和模拟可重复性，均不能替代外部工程验证。

5 讨论

5.1 预测、残差与反馈承担不同职责

基础模型将有限历史响应压缩为可计算的候选指标映射。残差层可以表示基础特征遗漏的非线性，但其有效性取决于残差样本量与偏差结构，不能仅凭“增加模型层”推断更好的阶段外预测。式(9)将预测误差与控制收益分开；二者之间还存在候选排序、约束和分位数目标的非线性作用。

反馈验收直接使用模拟器输出，因此能够拒绝预测目标较低但保护指标不合格的候选。有限库的完整枚举使最终反馈选择可被直接复算，同时也限制了残差消融的解释：本试验没有证明残差模型降低了模拟预算，也没有证明完整方法的收益由残差校正单独造成。当前结果更适合用于分析代理建议应如何

经过数值验收，以及小样本反馈的接受决定能否扩展到新工况。

5.2 适配预算与泛化边界

60+15+15 是本文的固定工况组织方式，不是从文献推导出的最优比例。既有候选库降低了新增阶段的搜索空间，因而 90 个适配工况的成本需要与库形成成本分开。后续可在固定总预算下改变三阶段配额，并以多个独立重复比较预测质量、可接受策略覆盖和测试结果；当前单次划分不足以确定最优配额。

候选在 15 个反馈工况上满足保护规则，只能支持该集合内的接受决定。面对不同初态、扰动或模型结构时，即使平均综合目标继续改善，个别保护指标也可能变化。因此，本文同时报告测试保护状态和连续效果，不以均值改善代替工程验收。辅助确认提供了更大工况集合上的特定策略结果，但两个实验不共享同一适配结论。

5.3 统计与工程接入

配对计算减少了同工况随机扰动对比较的干扰。bootstrap 以工况为重采样单位，保留基线与候选的对应关系；同一工况内的 121 个自相关时点不作为 121 个独立实验。不同模型族、固定工况分布及有限样本高分位估计仍限制了置信区间的解释。本文未证明联合显著性或分布无关保证。

努力度代理描述模拟执行器的相对代价，没有电压、电流或能耗测量支持，不能换算为实车节能率。匿名测点尚未与具体物理部件建立完整对应，参数敏感度也有待受控扰动试验识别。面向工程接入，需要补齐参数版本、测点字典和执行接口，再通过仿真、台架与实车逐级核验约束。云侧可保留候选说明、版本与验收结果，最终执行端仍应具备独立边界检查和基线回退能力。

6 结论

本文给出了一个基于有限策略库的 60+15+15 汽车热管理标定适配框架，明确区分基础指标拟合、低维残差校正、配对反馈选择与独立测试。优化在完整参数约束和初态路由规则内执行，采用逐工况归一化的均值-分位数目标，并保留逐项保护规则及完整基线。新适配试验中，完整反馈在六个上下文的测试保护通过数为 4，基础代理和残差代理分别为 3 和 2。两个兼具正收益与保护通过的反馈策略改善为 1.2360% 和 3.6809%，另有两个上下文保留基线。残差层未取得测试预测精度改善，且全候选枚举下不能将反馈收益归因于残差学习或模拟调用节约。

辅助确认表明，特定冻结分区策略可在所登记的 12 个模拟面板上取得 0.3650% 至 1.2659% 的综合目标改善，实际启用覆盖为 312/1440。研究支持将大模型建议和指标代理置于可复算的数值验收流程中；真实车辆效果、连续参数空间最优性和新增模型调用的独立收益仍需专门验证。

数据可用性声明 (Data Availability Statement)

原始运行记录和参数快照为脱敏工程材料，当前未公开发布。模拟工况、归一化参数、阶段切分及计算结果在研究环境中保存；经数据权属方许可后，可提供不包含原始参数基值与敏感测点的复现材料。本文未提供公开数据集或代码仓库的可用性承诺。

利益冲突声明 (Conflict of Interest Statement)

利益冲突披露将在署名作者确定后、正式提交前完成。

参考文献 (References)

- Charles Audet and J. E. Dennis, Jr. Mesh adaptive direct search algorithms for constrained optimization. *SIAM Journal on Optimization*, 17(1):188–217, 2006. doi: 10.1137/040603371. URL <https://doi.org/10.1137/040603371>.
- Andrew R. Conn, Katya Scheinberg, and Luis N. Vicente. *Introduction to Derivative-Free Optimization*. Society for Industrial and Applied Mathematics, 2009. doi: 10.1137/1.9780898718768. URL <https://epubs.siam.org/doi/book/10.1137/1.9780898718768>.
- B. Efron. Bootstrap methods: Another look at the jackknife. *The Annals of Statistics*, 7(1), 1979. doi: 10.1214/aos/1176344552. URL <https://doi.org/10.1214/aos/1176344552>.
- Arthur E. Hoerl and Robert W. Kennard. Ridge regression: Biased estimation for nonorthogonal problems. *Technometrics*, 12(1):55–67, 1970. doi: 10.1080/00401706.1970.10488634. URL <https://doi.org/10.1080/00401706.1970.10488634>.
- Qiuhaohu Hu, Mohammad Reza Amini, Ashley Wiese, Ilya Kolmanovsky, and Jing Sun. Robust model predictive control for enhanced fast charging on electric vehicles through integrated power and thermal management. arXiv:2310.13883, 2023. URL <https://arxiv.org/abs/2310.13883v1>.
- Marc C. Kennedy and Anthony O’Hagan. Bayesian calibration of computer models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 63(3):425–464, 2001. doi: 10.1111/1467-9868.00294. URL <https://doi.org/10.1111/1467-9868.00294>.
- Kyoung Hyun Kwak, Youyi Chen, Jaewoong Kim, Youngki Kim, and Dewey D. Jung. Thermal comfort-conscious eco-climate control for electric vehicles using model predictive control. *Control Engineering Practice*, 136:105527, 2023. doi: 10.1016/j.conengprac.2023.105527. URL <https://doi.org/10.1016/j.conengprac.2023.105527>.
- K. K. Gordon Lan and David L. DeMets. Discrete sequential boundaries for clinical trials. *Biometrika*, 70(3):659–663, 1983. doi: 10.1093/biomet/70.3.659. URL <https://academic.oup.com/biomet/article-abstract/70/3/659/247777>.
- Tennison Liu, Nicolás Astorga, Nabeel Seedat, and Mihaela van der Schaar. Large language models to enhance Bayesian optimization. In *International Conference on Learning Representations*, 2024. URL <https://arxiv.org/abs/2402.03921v2>.
- R. Tyrrell Rockafellar and Stanislav Uryasev. Optimization of conditional value-at-risk. *The Journal of Risk*, 2(3), 2000. doi: 10.21314/jor.2000.038. URL <https://www.risk.net/journal-risk/2161159/optimization-conditional-value-risk>.
- Wenyi Wang, Jiahang Ren, Xiang Yin, Yiyuan Qiao, and Feng Cao. Energy-efficient operation of the thermal management system in electric vehicles via integrated model predictive control. *Journal of Power Sources*, 603:234415, 2024. doi: 10.1016/j.jpowsour.2024.234415. URL <https://doi.org/10.1016/j.jpowsour.2024.234415>.
- Chengrun Yang, Xuezhi Wang, Yifeng Lu, Hanxiao Liu, Quoc V. Le, Denny Zhou, and Xinyun Chen. Large language models as optimizers. In *International Conference on Learning Representations*, 2024. URL <https://arxiv.org/abs/2309.03409v3>.