

When the Best Representation Changes: Permutation EDAs under Scalar and Pareto Selection

Anonymous Author(s)

Affiliation withheld for double-blind review

[VERIFY: institution, department, and contact details before submission]

Abstract—A permutation estimation-of-distribution algorithm (EDA) is defined by what its probabilistic model can encode — job adjacency, absolute position, or an explicit start position — and by how that model accumulates evidence: incrementally, with signed reward and punishment, or rebuilt each generation from a positive-only histogram. This paper asks whether the relative effectiveness of these representations changes when selection moves from scalar convergence to Pareto trade-off preservation, using permutation flow-shop scheduling (makespan and total flow time) as the test domain. We compare twelve EDA variants — six incremental COIN-family models and their punched-template ablations, and four rebuilt edge/node-histogram baselines — under matched population and evaluation budgets. Under bi-objective Pareto selection (Taillard TA001–TA030, thirty seeds, 10,800 runs), start-node/edge and hybrid node-edge representations rank highest on hypervolume and IGD+. Under scalar selection of either objective alone (TA001–TA020, four seeds, matched budget), punched-template histogram and template-grafted COIN variants rank highest instead, while the Pareto leaders fall to the middle of the field. The reversal is representation-specific, not template-specific: template inheritance itself helps the rebuilt edge histogram but hurts a rebuilt node histogram and both incremental COIN hosts it is grafted onto. We report this ranking reversal as a controlled empirical result, not as a state-of-the-art claim, and discuss it as evidence that objective structure, representation, and selection regime interact rather than that one representation dominates unconditionally.

Index Terms—estimation of distribution algorithms, permutation flow-shop scheduling, COIN algorithm, multi-objective optimization, hypervolume, IGD+, Pareto optimization, representation learning

I. INTRODUCTION

Two questions are routinely conflated in the estimation-of-distribution-algorithm (EDA) literature for permutation problems: *which probabilistic representation best captures a permutation problem’s structure*, and *which representation wins under the selection rule a particular study happened to use*. Most comparative studies fix the selection rule — almost always scalar convergence toward a single objective — and vary the representation. This leaves open whether an observed ranking is a property of the representation or an artifact of pairing that representation with scalar selection specifically.

This paper isolates the second factor. Using permutation flow-shop scheduling (PFSP) with makespan and total flow time as a matched pair of objectives, we run the same twelve EDA variants under two selection regimes — scalar minimization of each objective separately, and bi-objective

Pareto selection of both jointly — with identical population size, generation budget, and evaluation count in both regimes. If representation quality were selection-regime-independent, the same methods would rank similarly in both. They do not: the leading representations under Pareto selection are not the leading representations under scalar selection, and the reversal is not simply “templates always help” or “incremental learning always helps” — it is representation-specific.

We use the COIN family (an incremental, signed-evidence permutation EDA) and the EHBSA/NHBSA histogram family (rebuilt, positive-only-evidence EDAs) as the comparison, because the two families share representational targets in pairs — COIN and EHBSA both learn job-adjacency (edge) evidence; NB-COIN and NHBSA both learn absolute-position evidence — while differing in exactly one mechanism: whether evidence is accumulated incrementally with reward and punishment, or rebuilt from scratch each generation from selected individuals only. This pairing lets us separate the *evidence-accumulation* axis from the *representational-target* axis, and a further template-inheritance ablation (grafting punched-template resampling onto the incremental COIN models) lets us test whether template inheritance is a representation-general or representation-specific mechanism.

Central question: Does the relative effectiveness of a permutation EDA representation change when selection moves from scalar convergence to Pareto trade-off preservation?

This paper makes three contributions:

- 1) A controlled comparison of edge, position, start-node, and hybrid node-edge permutation EDA representations, with and without punched-template inheritance, under matched scalar and Pareto selection.
- 2) Empirical evidence of a representation-ranking reversal between scalar and Pareto selection, reported at the level of individual methods rather than as an aggregate family effect.
- 3) A cautious representation $\times$ objective-selection interpretation: no representation tested is uniformly superior, and template inheritance’s effect is host-dependent rather than universal.

II. RELATED WORK AND REPRESENTATION TAXONOMY

Permutation EDAs have been proposed for flow-shop scheduling under edge-histogram, node/position-histogram,

TABLE I
REPRESENTATION TAXONOMY

Method	Relation	Evidence rule	Templ.
COIN	Edge	Incr., signed	No
COIN/WT	Edge	Incr., signed	Yes
NB-COIN	Position	Incr., signed	No
NB-COIN/WT	Position	Incr., signed	Yes
CNB-COIN	Position (ord.)	Incr., signed	No
SNE-COIN	Start-node, edge	Incr., signed	No
HNE-COIN	Node templ., edge	Incr., signed	Partial
CNE-COIN	Chained node / edge	Incr., signed	Partial
EHBSA/WO	Edge	Rebuilt, pos.-only	No
EHBSA/WT	Edge	Rebuilt, pos.-only	Yes
NHBSA/WO	Position	Rebuilt, pos.-only	No
NHBSA/WT	Position	Rebuilt, pos.-only	Yes

and multi-objective simulated-annealing or memetic representations [6], [7], [2], [3], [4], [5], almost always evaluated under scalar makespan or total-flow-time minimization, or under multi-objective selection without a matched scalar re-run of the same representation. Multi-objective permutation EDA work exists but has rarely re-examined the same representations that were validated under scalar selection; a representation’s scalar-regime ranking is typically inherited into multi-objective settings without re-testing. COIN, the incremental signed-evidence family used here, has itself been evaluated almost exclusively under scalar objectives in its own prior applications to PFSP — node-based COIN for total flow time [8] and, separately, for makespan [9]. This paper’s gap is specifically the absence of a same-implementation, matched-budget comparison across both selection regimes.

Before introducing the twelve variants, two structural distinctions matter. First, *learned relation*: an edge-based model (COIN, EHBSA) learns which job follows which; a position-based model (NB-COIN, NHBSA) learns which job occupies which absolute position — a distinction already recognized outside PFSP [18] and consistent with classical arguments that permutation encodings need order- or adjacency-based building blocks rather than fixed-locus schemata [19], [20]; a start-node model (SNE-COIN) additionally learns an explicit first-job distribution, which edge- and position-only models do not represent as a separate event. Second, *evidence accumulation*: COIN-family models update an existing weight matrix incrementally, adding signed reward and punishment each generation [15], [17]; EHBSA/NHBSA-family models discard their matrix each generation and rebuild it as a positive-only histogram from the current elite selection only.

Research question: Does the ranking induced by this taxonomy under scalar selection match the ranking it induces under Pareto selection?

III. METHODS AND EXPERIMENTAL PROTOCOL

All twelve variants share a population of 100 and a training rate of 5. COIN-family variants apply reward to the top 5%

and punishment to the bottom 5% of the ranked population each generation, updating their weight matrix incrementally rather than replacing it; EHBSA/NHBSA-family variants rebuild their histogram from the top 50% (elite selection) each generation, with a bias ratio of 0.005. Template variants (EHBSA/WT, NHBSA/WT, COIN/WT, NB-COIN/WT) additionally punch and resample roughly 50% of a parent permutation with one-to-one replacement before the representation-specific sampling step runs on the remaining positions; HNE-COIN and CNE-COIN use a related but structurally distinct template mechanism (a retained/chained node scaffold rather than a punched-and-resampled parent) and are kept as separate rows rather than folded into the WT ablation family.

For each edge- or position-histogram model, the reward–punishment identity is: COIN accumulates signed evidence toward H_{ij} (edge) or $H_{k,j}$ (position) from both rewarded and punished individuals across generations; EHBSA/NHBSA discard H and rebuild it from $10C_{ij}^+ + 1$ (or the position equivalent) using only the current generation’s elite cohort, with no memory of prior generations.

Instances and budget. Taillard PFSP instances TA001–TA020 (the twenty instances common to every method under both regimes). Population 100, 400 generations, 40,000 evaluations per run, identical across regimes and objectives [1].

Scalar regime. Each of makespan and total flow time is optimized separately as a single scalar objective. Four paired seeds (42, 2026, 2027, 2028) $\times$ 20 instances $\times$ 12 methods = 960 runs per objective, 1,920 runs total. Performance is reported as ARPD (average relative percentage deviation from the best value observed for that instance across all compared methods) and mean per-block rank, averaged over instance–seed blocks.

Pareto regime. Makespan and total flow time are optimized jointly; selection uses Pareto depth plus a crowding tiebreak (no scalarization). Thirty seeds $\times$ 30 instances $\times$ 12 methods = 10,800 runs, with no recorded run failures. Hypervolume (HV) is computed after a fixed per-instance min–max normalization, reference point (1.1, 1.1) [10]; IGD+ (not ordinary IGD) is computed against a pooled observed nondominated front built from all twelve methods’ complete seeds for that instance [11]. A Friedman test over instance $\times$ seed blocks ($n = 900$) tests the omnibus null of equal performance [12]; Holm-corrected pairwise Wilcoxon tests compare the leading method against every competitor [13], [14].

Both regimes checked output validity (complete, feasible tours) before recording a run; no additional correctness verification (e.g., cross-implementation equality) is claimed for this comparison specifically.

IV. SCALAR RESULTS

The template-inheriting methods dominate scalar selection. NHBSA/WT ranks first on makespan (ARPD 0.334%) and second on total flow time (0.542%); NB-COIN/WT ranks second on makespan (0.482%) and first on total flow time (0.315%). Among non-template methods, NB-COIN is the strongest scalar performer (rank 3 makespan, though only rank

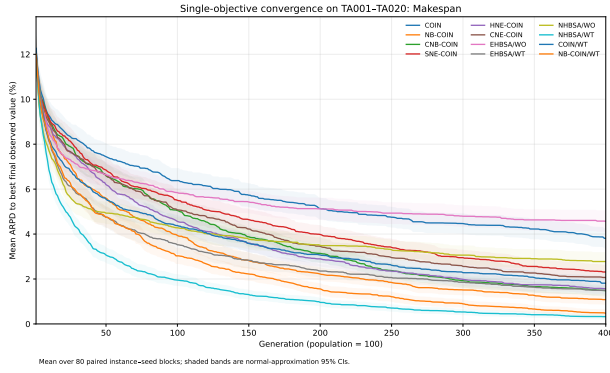


Fig. 1. (a) Mean ARPD (%) vs. generation for scalar makespan minimization, TA001–TA020, four seeds, 95% confidence bands. Template-inheriting methods separate from non-template methods within the first 100 generations and remain separated at convergence.

7 total flow time), while COIN and EHBSA/WT occupy the last two positions on both objectives (COIN: 3.821%/5.329%; EHBSA/WT: 4.575%/6.773% ARPD) — the incremental edge model without a template and the rebuilt edge histogram without a template are, respectively, the weakest incremental and weakest rebuilt method under scalar selection.

Fig. 1 shows mean ARPD versus generation for scalar makespan minimization, reused from the source benchmark run; the template-inheriting curves separate from the non-template curves within roughly the first 100 generations and remain separated at convergence, consistent with template inheritance providing a persistent rather than transient scalar advantage for the two histogram-family hosts (EHBSA, NHBSA) and the two COIN hosts it was grafted onto for this specific ablation.

This scalar comparison rests on four seeds over twenty instances (80 instance–seed blocks per objective) — a matched, paired-seed design, but a smaller evidence base than the Pareto regime’s thirty seeds. We treat scalar rankings as directionally reliable but do not apply significance testing calibrated for a larger sample to this four-seed comparison; Section VI states this explicitly as a limitation rather than silently equalizing the two regimes’ evidentiary weight.

V. PARETO RESULTS

Under Pareto selection, the ranking inverts by representation. SNE-COIN attains the best mean hypervolume (0.9871, reference point (1.1, 1.1)) and the third-best mean IGD+ (0.1149); HNE-COIN is second on both indicators (HV 0.9808, IGD+ 0.1117); NB-COIN attains the best mean IGD+ (0.1096) while ranking third on HV (0.9708). A Friedman test over the 900 instance–seed blocks rejects equal performance across all twelve methods ($\chi^2 = 8040.47$, $p < 10^{-300}$); Holm-corrected pairwise comparison of the HV leader (SNE-COIN) against every competitor is significant throughout, including against its closest competitor HNE-COIN (median $\Delta HV = 0.0186$, $p_{\text{Holm}} = 0.000183$) and most strongly against the weakest method, COIN/WT (median $\Delta HV = 0.3724$, $p_{\text{Holm}} \approx 7.4 \times 10^{-148}$).

TABLE II

COMBINED REGIME COMPARISON: ORDINAL RANK (1 = BEST) OF EACH METHOD UNDER SCALAR MAKESPAN, SCALAR TOTAL FLOW TIME (TFT), PARETO HYPERVOLUME (HV), AND PARETO IGD+. SCALAR RANKS ARE RECALCULATED FOR THIS TWELVE-METHOD MATCHED SUBSET (ROSE/WT EXCLUDED, SEE TEXT).

Method	Sc.-MS	Sc.-TFT	HV	IGD+
SNE-COIN	9	9	1	3
HNE-COIN	6	6	2	2
NB-COIN	3	7	3	1
CNB-COIN	4	4	4	4
CNE-COIN	8	8	5	5
NHBSA/WT	10	10	6	6
NHBSA/WT	1	2	7	7
EHBSA/WT	5	3	8	8
COIN	11	11	9	9
NB-COIN/WT	2	1	10	10
EHBSA/WT	12	12	11	11
COIN/WT	7	5	12	12

The two methods that were strongest under scalar selection are, by contrast, mid-to-weak under Pareto selection: NHBSA/WT falls to rank 7 of 12 on both HV and IGD+ (HV 0.8866, IGD+ 0.1865), and NB-COIN/WT falls to rank 10 of 12 on both (HV 0.6996, IGD+ 0.3118). COIN/WT — the template grafted directly onto the plain incremental edge model — is the single worst method under Pareto selection (HV 0.6205, rank 12 of 12), despite ranking mid-field (7th of 12 on makespan, 5th on total flow time) under scalar selection. This is the clearest single-method instance of the reversal: template inheritance’s benefit is not portable from its native histogram-rebuild hosts to an incremental host, and under Pareto selection that same graft is actively harmful rather than neutral.

Table II makes the reversal pattern visible at a glance: eight of the twelve methods move by four or more rank positions between the scalar-makespan and Pareto-HV columns; only CNB-COIN (ranks 4th under every one of the four columns) and, to a lesser extent, NHBSA/WT (ranks 6th on both MO indicators, 10th on both scalar objectives — consistently weak rather than consistently strong) hold a stable position across regimes.

Fig. 2 plots each method’s scalar-makespan rank against its Pareto mean-HV rank as a slope graph; the crossing lines visualize the reversal more directly than the table alone. Five methods — NHBSA/WT, NB-COIN/WT, SNE-COIN, HNE-COIN, and COIN/WT — move by six or more rank positions and are highlighted.

VI. DISCUSSION, LIMITATIONS, AND CONCLUSION

A. Discussion

The safe reading of this result is representation-specific, not family-general: a representation that supports rapid concentration under scalar selection need not be the representation that best supports Pareto-front approximation, and this is true within each representational family, not only between families. We deliberately avoid the stronger claim that Pareto selection

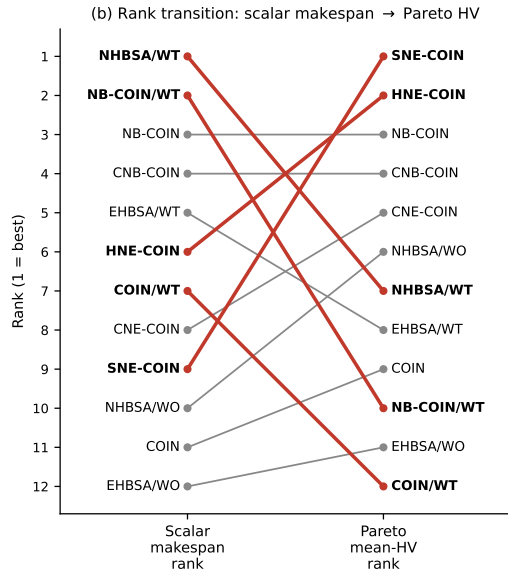


Fig. 2. (b) Rank-transition slope graph: each method’s scalar-makespan rank (left) against its Pareto mean-hypervolume rank (right). Highlighted lines mark a representation whose scalar and Pareto rankings disagree by six or more positions.

directly rewards structurally different permutations — Pareto selection rewards objective-space trade-offs, and structural heterogeneity in the resulting population is a plausible mechanism, not one measured here. COIN’s signed, incremental evidence retaining both rewarded and punished information is consistent with a hypothesis that it preserves more structural diversity than a rebuilt positive-only histogram, but we have not measured model entropy, edge/position entropy, or pairwise permutation distance in this study, so we present this as a hypothesis the reversal is consistent with, not a proven mechanism.

B. Limitations

(1) The scalar and Pareto regimes carry different statistical weight — four seeds versus thirty — so scalar rankings should be read as directionally suggestive rather than confirmatory at the same standard as the Pareto results; a matched 20+-seed scalar re-run is the natural next step, not attempted here. (2) IGD+ is computed against a pooled observed reference front built from the same twelve methods being compared, which can favor methods that are more typical of the pool; an external or hypervolume-only cross-check partially mitigates this but does not eliminate it. (3) We exclude general-purpose GA and NSGA/SPEA2 baselines by design — the research question is about EDA representation, not state-of-the-art dominance over non-EDA methods — so this study does not position any tested method against the wider metaheuristic literature. (4) We report no decision-space diversity or model-entropy measurement, so the structural-diversity explanation for the reversal remains a hypothesis for future instrumented study, not a finding of this paper.

C. Conclusion

Twelve permutation EDA variants, sharing representational targets in matched pairs and evaluated under identical population and evaluation budgets, rank differently under scalar and Pareto selection, and the reversal is method-specific rather than a single family effect: punched-template inheritance helps its two histogram-family hosts under scalar selection, but is neutral-to-harmful for solution diversity under Pareto selection, most severely so when grafted onto an incremental edge model. No representation tested is uniformly superior across both regimes. This motivates evaluating a permutation EDA’s representation under the selection regime it will actually be deployed with, rather than assuming a scalar-regime ranking transfers to a multi-objective setting.

DATA AND CODE AVAILABILITY

The anonymized submission will provide configurations, seeds, raw archives, and analysis scripts, subject to the venue’s double-anonymous policy.

REFERENCES

- [1] E. Taillard, “Benchmarks for basic scheduling problems,” *European Journal of Operational Research*, vol. 64, no. 2, pp. 278–285, 1993. doi:10.1016/0377-2217(93)90182-M.
- [2] T. K. Varadharajan and C. Rajendran, “A multi-objective simulated-annealing algorithm for scheduling in flowshops to minimize makespan and total flowtime,” *European Journal of Operational Research*, vol. 167, no. 3, pp. 772–795, 2005. doi:10.1016/j.ejor.2004.07.020.
- [3] G. Minella, R. Ruiz, and M. Ciavotta, “A review and evaluation of multiobjective algorithms for the flowshop scheduling problem,” *INFORMS Journal on Computing*, vol. 20, no. 3, pp. 451–471, 2008. doi:10.1287/ijoc.1070.0258.
- [4] G. Minella, R. Ruiz, and M. Ciavotta, “Restarted iterated Pareto greedy algorithm for multi-objective flowshop scheduling problems,” *Computers & Operations Research*, vol. 38, no. 11, pp. 1521–1533, 2011. doi:10.1016/j.cor.2011.01.010.
- [5] T.-C. Chiang, H.-C. Cheng, and L.-C. Fu, “NNMA: An effective memetic algorithm for solving multiobjective permutation flow shop scheduling problems,” *Expert Systems with Applications*, vol. 38, no. 5, pp. 5986–5999, 2011. doi:10.1016/j.eswa.2010.11.022.
- [6] S. Tsutsui and M. Miki, “Solving flow shop scheduling problems with probabilistic model-building genetic algorithms using edge histograms,” in *Proc. 4th Asia-Pacific Conf. Simulated Evolution and Learning (SEAL 2002)*, 2002, pp. 776–780. doi:10.1142/9789812561794_0013.
- [7] S. Tsutsui, M. Pelikan, and D. E. Goldberg, “Node histogram vs. edge histogram: A comparison of probabilistic model-building genetic algorithms in permutation domains,” in *Proc. IEEE Congress on Evolutionary Computation*, 2006, pp. 1939–1946. doi:10.1109/CEC.2006.1688444. [VERIFY once more against IEEE Xplore directly; secondary indexes can corrupt this DOI]
- [8] O. Srimongkolkul and P. Chongstitvatana, “Application of Node Based Coincidence algorithm for flow shop scheduling problems,” in *Proc. 10th Int. Joint Conf. Computer Science and Software Engineering (JCSSE)*, 2013, pp. 49–52. doi:10.1109/JCSSE.2013.6567318.
- [9] O. Srimongkolkul and P. Chongstitvatana, “Minimizing makespan using node based coincidence algorithm in the permutation flowshop scheduling problem,” in *Industrial Engineering, Management Science and Applications 2015*, M. Gen et al., Eds., Lecture Notes in Electrical Engineering, Springer, 2015, pp. 303–311. doi:10.1007/978-3-662-47200-2_33.
- [10] A. P. Guerreiro, C. M. Fonseca, and L. Paquete, “The hypervolume indicator: Computational problems and algorithms,” *ACM Computing Surveys*, vol. 54, no. 6, 2021. [VERIFY: article number/DOI]
- [11] H. Ishibuchi, H. Masuda, Y. Tanigaki, and Y. Nojima, “Modified distance calculation in generational distance and inverted generational distance,” in *Evolutionary Multi-Criterion Optimization*, 2015. [VERIFY: LNCS pages/DOI]

- [12] M. Friedman, "The use of ranks to avoid the assumption of normality implicit in the analysis of variance," *Journal of the American Statistical Association*, vol. 32, no. 200, pp. 675–701, 1937.
- [13] S. Holm, "A simple sequentially rejective multiple test procedure," *Scandinavian Journal of Statistics*, vol. 6, no. 2, pp. 65–70, 1979.
- [14] J. Demšar, "Statistical comparisons of classifiers over multiple data sets," *Journal of Machine Learning Research*, vol. 7, pp. 1–30, 2006.
- [15] W. Wattanapornprom, P. Olanvitchai, P. Chutima, and P. Chongstitvatana, "Multi-objective combinatorial optimisation with coincidence algorithm," in *Proc. IEEE Congress on Evolutionary Computation (CEC)*, 2009, pp. 1675–1682. doi:10.1109/CEC.2009.4983143.
- [16] W. Wattanapornprom and P. Chongstitvatana, "Solving multimodal combinatorial puzzles with edge-based estimation of distribution algorithm," in *Proc. Genetic and Evolutionary Computation Conference Companion (GECCO'11 Companion)*, Dublin, Ireland, 2011. **[VERIFY: page numbers and DOI from the ACM Digital Library record before submission]**
- [17] W. Wattanapornprom and P. Chongstitvatana, "Negative correlation learning in the estimation of distribution algorithms for combinatorial optimization," *IEICE Transactions on Information and Systems*, vol. E96-D, no. 11, pp. 2397–2408, 2013. doi:10.1587/transinf.E96.D.2397.
- [18] K. Waiyapara, W. Wattanapornprom, and P. Chongstitvatana, "Solving Sudoku puzzles with node based Coincidence algorithm," in *Proc. 10th Int. Joint Conf. Computer Science and Software Engineering (JCSSE)*, 2013. doi:10.1109/JCSSE.2013.6567311.
- [19] D. E. Goldberg and R. Lingle Jr., "Alleles, loci, and the traveling salesman problem," in *Proc. 1st Int. Conf. Genetic Algorithms and Their Applications (ICGA)*, 1985, pp. 154–159. **[VERIFY: the primary 1985 proceedings text was not directly read for this manuscript; the characterization given in Section II is a widely repeated secondary characterization, not a direct quotation]**
- [20] P. Larrañaga, C. M. H. Kuijpers, R. H. Murga, I. Inza, and S. Dizdarevic, "Genetic algorithms for the travelling salesman problem: A review of representations and operators," *Artificial Intelligence Review*, vol. 13, no. 2, pp. 129–170, 1999. doi:10.1023/A:1006529012972.