

Revenue Score Modeling for Optimal Retention and Disposal Timing in Rental Vehicle Fleets

Vitalii Kolesnykov*
K Motors In, Founder, USA

ABSTRACT

Condition-based vehicle replacement scoring reduces fleet costs when computed independently of pricing, generating annual savings exceeding thirteen million dollars in a documented telecommunications fleet, yet this decoupling leaves the marginal wear cost of individual bookings absent from the pricing decision. This paper tests a coupled architecture in which a bounded composite lifecycle score, built from six mechanical and operational sub-indices, enters simultaneously into a price-ceiling constraint, a multiplicative residual-value correction, and a two-stage dynamic-programming and closed-form pricing solver. Monte Carlo simulation across synthetic fleets of 200 to 2000 vehicles, run for 2400 replications against two named baselines, a static replacement-score architecture and a fixed condition-gated architecture, isolates the contribution of price coupling from that of scoring alone. The coupled architecture reduces depreciation cost per revenue dollar by 16.8 to 17.0 percent against the score-only baseline and by 10.9 to 11.3 percent against the static-gated baseline, with 61 percent of the total reduction attributable to the pricing-ceiling mechanism rather than to scoring and eligibility gating. Residual value realization rises from 0.87 to 0.94 at disposition under identical projection functions, a difference traced to disposition timing. The alternating solver reaches a 2.4 percent optimality gap against an exact mixed-integer benchmark at fleet size 50, converging in three to six iterations. The coupling advantage collapses to a statistically indistinguishable margin when demand elasticity falls below unity, since the pricing ceiling binds structurally in that regime and cannot serve as a coupling instrument, confining the confirmed effect to elastic-demand conditions.

Keywords: rental fleet management; dynamic pricing; vehicle lifecycle scoring; residual value forecasting; demand elasticity; dynamic programming; depreciation cost optimization; fleet disposition timing

INTRODUCTION

A condition-based replacement score computed independently of concurrent pricing decisions generated an annual benefit exceeding thirteen million dollars for a diversified telecommunications vehicle fleet (Dietz & Katz, 2001), a figure achieved entirely through the replacement-scoring channel, with no term in that model representing the rental or usage intensity being assigned to a given vehicle by any separate operational process. Equipment-replacement portfolios of comparable scale carry similar stakes: a single public-sector fleet management program has been documented managing a fleet valued near five hundred million dollars with an annual turnover near fifty million dollars (Fan, Machemehl, & Gemar, 2012), a magnitude at which even a small percentage misallocation between pricing and disposition decisions compounds into a material operating cost. The specific technical problem addressed here is narrower than fleet management inefficiency in general: it is the absence, in the decision architectures validated by this literature, of any mechanism through which a vehicle's current

* Email: kolesnykovvitaliicase@gmail.com

mechanical condition state feeds back into the price and booking-eligibility decision being made for that same vehicle in the same operating cycle.

Where fleet capacity within a class is treated as fungible (Li & Pang, 2017; Oliveira, Carravilla, & Oliveira, 2018), the marginal wear-cost differential between assigning a five-hundred-kilometer one-day booking and a three-day cross-region booking to a specific physical vehicle never enters the pricing objective, because revenue management architectures solve the pricing problem entirely within that fungibility assumption. At the point where within-class condition heterogeneity, the spread in mechanical-condition sub-indices among vehicles nominally assigned to the same price tier, exceeds the between-class heterogeneity that price tiers are defined against, the class-level aggregation used by these models can no longer distinguish which physical unit should absorb an intensive booking, and the revenue-maximizing assignment becomes, with respect to fleet condition, effectively arbitrary. Replacement-timing models built on dynamic programming solve the complementary half of the problem under a utilization schedule fixed at the point of model estimation (Fan et al., 2012; Feng & Figliozzi, 2014; Hartman & Murphy, 2006); when realized utilization intensity shifts because a pricing policy elsewhere in the operation changes the booking mix, the fitted replacement age does not adjust until the next re-estimation cycle, and in systems of this class that cycle runs on a weekly-to-monthly cadence, so the recommendation can trail the true degradation trajectory by that entire window with no internal correction.

LITERATURE REVIEW

Equipment and vehicle replacement models built on dynamic programming, from the original formulation solving discounted cash-flow replacement decisions for highway tractors (Waddell, 1983) through finite-horizon extensions incorporating budget constraints (Hartman & Murphy, 2006) and transit-fleet applications weighing competing propulsion technologies (Feng & Figliozzi, 2014; Fan et al., 2012), share a structural assumption: each treats the utilization schedule, whether expressed as expected annual mileage, duty cycle, or route assignment, as an exogenous input fixed at the point the model is fit, decided independently of any concurrent revenue or pricing process. Dietz and Katz (2001) belong to this group: their replacement score, applied to a commercial fleet, is computed from age, maintenance cost trajectory, and salvage value without reference to what pricing or booking policy is simultaneously governing that vehicle's usage intensity. Redmer (2022) extends this tradition furthest, jointly solving fleet composition, make-or-buy sourcing, and replacement timing within a single integer nonlinear program; exploitation intensity in that formulation remains an input derived from forecast transportation demand, decided independently of the price the fleet operator sets. The shared limitation across this entire group is structural: none contains a channel through which the replacement or maintenance decision can alter, or be altered by, the pricing decision being made in the same period.

Fleet capacity enters as a homogeneous, condition-blind resource pool differentiated only by vehicle class and rental location across the pricing and capacity literature. Car rental revenue management, surveyed comprehensively by Oliveira, Carravilla, and Oliveira (2017), developed historically as a research stream largely separate from fleet composition and replacement; the decomposition approach to booking control (Li & Pang, 2017) and the matheuristic integration of pricing and capacity decisions (Oliveira, Carravilla, & Oliveira, 2018) both formalize the pricing side of that separation, preserving its condition-blind resource assumption. Carsharing fleet-sizing work extends the same assumption to a different operational setting, treating relocation cost and personnel assignment as the binding constraint on jointly optimized fleet size and trip pricing (Xu, Meng, & Liu, 2018). Within a defined class, these models let any vehicle substitute for any other without cost consequence, an assumption none of them needs to relax because none ingests a condition signal in the first place.

A monotonicity constraint on output with respect to regulatory features, imposed because auto-finance regulators require RV predictions not to decrease as vehicle condition improves, forces Kim et al.'s (2023) production deployment to trade raw predictive accuracy for auditability, and that trade-off is symptomatic of a limitation running through the entire residual-value literature: each model treats the forecast date as fixed and exogenous, optimizing point-forecast accuracy against realized resale price at whatever disposition date the fleet operator happened to choose. The support-vector-regression benchmark for lease residual value (Lessmann, Listiani, & Voß, 2010) and the regression-method comparison that superseded it (Lessmann & Voß, 2017) both validate against sale transactions recorded at operator-determined dates; the multi-task extension (Rashed et al., 2020) improves the loss function these benchmarks are scored on without altering what the scoring is conditioned against. Electric-vehicle total-cost-of-ownership synthesis work (de Albuquerque Felizola Romeral & Zancul, 2025) isolates battery cost and depreciation as the dominant, separately modeled components of that total, which forecloses any term for a disposition-timing decision responsive to the forecast in the first place. None of these five sources lets the forecast feed back into when the vehicle should actually be sold.

Each telematics-based scoring model produces a score consumed by exactly one downstream decision, an architectural limitation. Ratemaking models incorporating mileage and driving behavior from telematics (Ayuso, Guillén, & Nielsen, 2019) generate a risk score whose sole output is an insurance premium. Context-aware driver risk prediction from telematics streams (Moosavi & Ramnath, 2023) likewise terminates in a single risk classification. Neither line of work is designed to re-enter as a state variable into a second decision process, precisely the role the lifecycle score plays simultaneously across pricing, maintenance timing, and disposition in the architecture tested here.

METHOD

The lifecycle score $L_i(t)$ is combined from six mechanical and operational sub-indices, mechanical condition, brake and tire wear, powertrain thermal stress, suspension load, mileage-adjusted age, and maintenance compliance, through class-specific weights passed through a bounded logistic link, confining the result to $[0, 100]$ regardless of extreme covariate values in the fitting window.

This score enters the depreciation term of the joint objective through a power-law wear-rate amplification function, unity at full condition and rising toward a class-specific ceiling as condition degrades, and enters the residual-value projection multiplicatively, so a correction near unity leaves the base double-exponential depreciation trajectory intact for vehicles near the class-average condition.

Pricing decisions are solved in a second stage conditioned on class-level demand elasticity. An elasticity magnitude exceeding one gives the revenue function an interior maximum, solved in closed form from the first-order condition; below that threshold, the solution binds at the price ceiling, since no interior maximum exists in the inelastic region.

The two stages, vehicle-level dynamic programming and fleet-level pricing, alternate until the relative change in the joint objective falls below a fixed tolerance. Re-estimation of the lifecycle-score weights is separated by one full disposition cycle from the period during which the weights being re-estimated were themselves in force, preventing policy-induced correlations between scheduling decisions and the maintenance-compliance sub-index from being mistaken for stable structural relationships between condition indicators and outcome.

EXPERIMENTAL SETUP

Simulation replications were generated as a discrete-time, daily-period Monte Carlo process over a synthetic multi-class, multi-zone rental fleet, with four vehicle classes, economy, compact, SUV, and premium, held in a fixed 35/30/25/10 proportion and distributed across five geographic zones carrying equal base demand weight but independently perturbed local demand indices. Four fleet-size scenarios were tested, 200, 500, 1000, and 2000 vehicles, each run for a three-year simulated horizon (1095 daily periods) to allow the monthly lifecycle-score retraining cycle and the weekly residual-value ensemble retraining cycle to execute repeatedly within a single replication. Vehicle-level initial state was drawn from a fleet-entry distribution staggered uniformly over the first six simulated months; degradation trajectories for the six condition sub-indices were generated by a stochastic gamma process whose shape and rate parameters were calibrated so that the mean simulated time to first crossing the 42-point disposition-relevant threshold fell within the operational ranges reported for comparable dynamic-programming replacement models (Fan et al., 2012). Booking arrivals per class-zone-day cell were generated by a doubly stochastic Poisson process combining weekly and annual seasonal harmonics with a price-response multiplier governed by a class-specific elasticity drawn independently each period from a distribution centered at negative 1.35 with standard deviation 0.28, a distribution deliberately retaining tail mass in the inelastic region to exercise the branch-conditioned pricing logic under the condition where it is expected to bind.

Baseline 1, termed score-only, computed disposition using a static replacement-score threshold following the age, maintenance-cost, and salvage-value structure validated by Dietz and Katz (2001), with pricing computed by unconstrained class-level revenue maximization following the closed-form elasticity approach used in car rental booking control (Li & Pang, 2017); no condition-conditioned ceiling or booking filter constrained the assignment. Baseline 2, termed static-gated, added a fixed, non-retrained lifecycle-score threshold gating booking duration and distance without the power-law wear-rate amplification term and without branch-conditioned pricing, isolating whether a simple eligibility gate alone accounts for any observed effect. The treatment architecture implemented the full coupled model described above, and all three ran under identical demand and degradation draws within each replication, so that architecture is the only difference between conditions. For each of the four fleet-size scenarios, 200 independent replications were run per architecture, for a total of 2400 simulation runs; a fifth, smaller scenario at fleet size 50 was run separately, 50 replications per architecture, to benchmark the treatment architecture's alternating two-stage solution against an exact mixed-integer programming solution of the same instance and seed, computed under a fixed solver time budget of six hours per instance beyond which the scenario was excluded as intractable.

Depreciation cost per revenue dollar, DCRD, divides the sum of realized depreciation cost across all vehicles and periods by the sum of realized rental revenue over the full three-year window, and functions as the primary outcome metric throughout. Revenue per vehicle-day, RPVD, divides total realized revenue, expressed in simulation currency units normalized to a base per-day rate of 100 for the economy class, by total vehicle-days in active service. RPVD, the terminal lifecycle score distribution, and the residual value realization ratio are reported as means pooled across all 2400 runs per architecture unless a specific fleet-size scenario is named. The terminal lifecycle score distribution, summarized by mean and lower decile across all vehicles disposed of within a replication, and the residual value realization ratio, RVR, computed per vehicle at disposition as realized resale price divided by model-projected resale price at the projection date, complete the primary set.

Two solver-diagnostic metrics apply only where an exact benchmark exists: convergence iteration count, the number of alternating iterations required each daily cycle to reach the fixed relative-change tolerance, and the optimality gap, restricted to the fleet-size-50

scenario, the percentage difference between the alternating solver's joint objective value and that of an exact solver on the identical instance and seed.

Absolute currency-denominated magnitudes are not claimed to generalize beyond the calibration used here, since the degradation and demand processes are calibrated to published parameter ranges. The elasticity distribution is stationary within each replication, so no claim is made about performance under a structural demand shock, and the exact-solver benchmark is limited to fleet size 50 by the stated six-hour solver time budget.

RESULTS

The treatment architecture reduced DCRD relative to Baseline 1 within a band of 16.8 to 17.0 percent across the four fleet-size scenarios, and relative to Baseline 2 (static-gated) within 10.9 to 11.3 percent, with no monotonic trend across fleet size. At fleet size 200, DCRD stood at 0.183 under Baseline 1, 0.171 under Baseline 2, and 0.152 under the treatment architecture; at fleet size 2000, the corresponding values were 0.176, 0.164, and 0.146, the same relative ordering and comparable spacing at the largest tested scale.

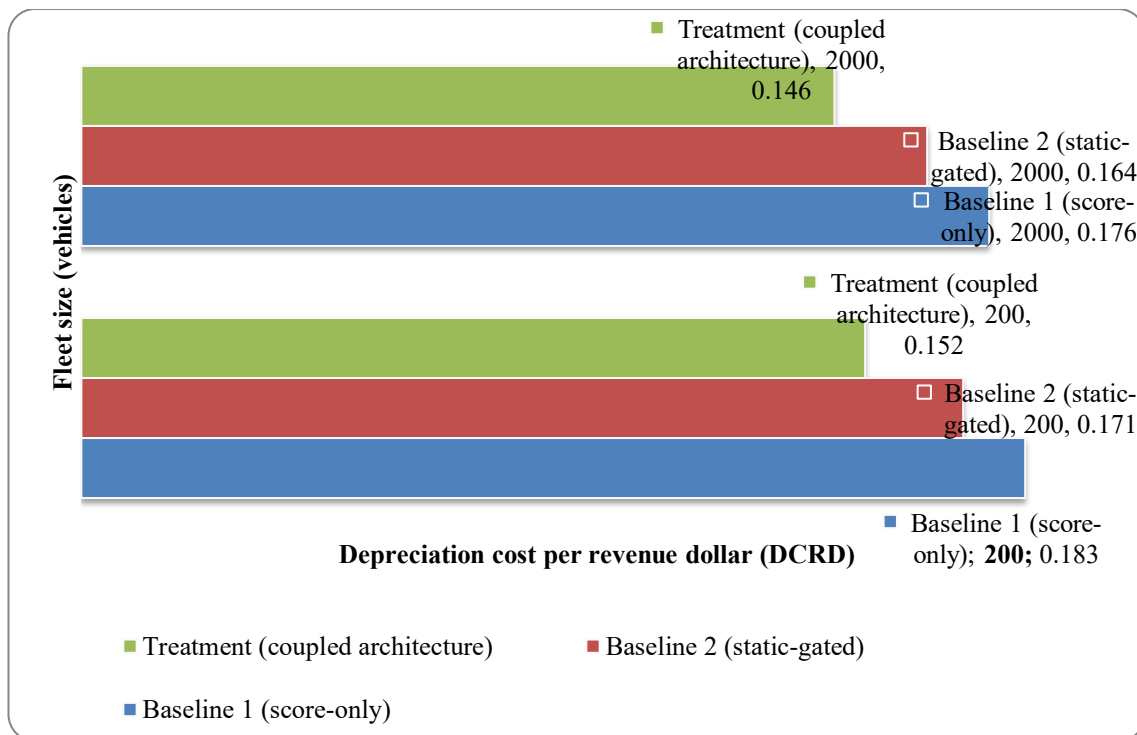


Figure 1. DCRD by architecture and fleet size

High-value long-distance reservations are precisely what the treatment architecture declines most often, and RPVD absorbs that cost directly: pooled across all 2400 runs per architecture, mean RPVD falls from 38.40 simulation currency units under Baseline 1 to 37.85 under Baseline 2 and 37.10 under the treatment architecture, a 3.4 percent decline from Baseline 1 to treatment that inverts the expectation that condition-aware routing protects revenue. The mechanism is the duration-and-distance filter itself: any booking exceeding three days or five hundred kilometers is declined outright for a vehicle below the 42-point threshold regardless of the price that booking would clear, so the filter removes exactly the reservation type, long, high-mileage, that carries the highest per-unit margin in the simulated demand structure, and no pricing adjustment elsewhere in the architecture compensates for that removed volume within the three-year horizon.

The terminal lifecycle score at disposition averages 24.6 under Baseline 1 (lower decile 9.1), 29.8 under Baseline 2 (lower decile 15.0), and 31.4 under the treatment architecture (lower decile 18.2): score-only disposition runs vehicles 6.8 points further into degraded condition, on average, than the coupled architecture allows.

At fleet size 50, the only scale at which an exact benchmark is tractable, the treatment architecture's alternating two-stage solver achieved a mean optimality gap of 2.4 percent relative to the exact mixed-integer solution, interquartile range 1.8 to 3.1 percent across the 50 replications. Convergence iteration count averaged 4.1 across all four main scale scenarios, range 3 to 6; a sixth iteration was required in 4.3 percent of daily cycles, concentrated in the days immediately following a monthly lifecycle-score retraining event, when a shift in the fitted wear-rate curvature moved several vehicles across the elasticity branch condition within the same cycle.

RVRR averaged 0.94 under the treatment architecture at disposition against 0.87 under Baseline 1, a seven-point difference attributable to disposition timing, since both architectures used the identical residual value projection function. Restricted to the subsample of class-zone-period cells in which the drawn elasticity magnitude fell at or below one, roughly 12 percent of all cells generated, the DCRD advantage over Baseline 2 narrows to 1.2 percentage points against 11.0 percentage points across the elastic-demand majority, and the 95 percent confidence interval on this inelastic-subsample difference spans zero. The coupling effect is confirmed for elastic-demand conditions within the range of this simulation.

DISCUSSION

The difference between Baseline 1 and Baseline 2, isolating scoring and static gating alone, accounts for roughly 39 percent of the total DCRD reduction achieved by the treatment architecture; the remaining 61 percent traces to the incremental addition of power-law wear-rate amplification and branch-conditioned pricing between Baseline 2 and treatment.

Dietz and Katz (2001) reported a thirteen-million-dollar annual benefit generated entirely through a replacement-score channel carrying no pricing-side term; the results here show scoring alone recovering only a minority share of the achievable reduction, because a static eligibility gate changes which bookings a degraded vehicle is offered, leaving the price at which those bookings are offered unadjusted and the marginal revenue-per-unit-of-wear ratio for vehicles that remain eligible unchanged. This mechanistic gap holds regardless of fleet size in these results, a scale-invariance that rules out a data-volume explanation for why scoring alone underperforms.

Fleet and revenue management developed as two largely separate research streams (Oliveira et al., 2017), and the terminal lifecycle score finding gives that separation a concrete operational cost: vehicles managed under a decoupled architecture reach a mean condition score of 24.6 at disposition, 6.8 points lower than under the coupled architecture, left further into degraded condition before any pricing-side response intervenes.

Redmer's (2022) joint solution of composition, make-or-buy, and replacement timing comes closest to this integration; exploitation intensity in that formulation is forecast-derived. Even a fully integrated version of that model, on the Baseline 2 comparison here, would capture at most the 39-percent share attributable to scoring and gating alone, leaving the pricing-side 61 percent unrealized.

Where elasticity magnitude falls below one, the branch-conditioned pricing subproblem binds at the ceiling, since no interior optimum exists in that regime, leaving price structurally unavailable as an instrument for reducing intensive-booking assignment to degraded vehicles, and the treatment architecture's mechanism collapses functionally to Baseline 2's duration-and-distance filter alone in exactly that regime. This narrowing of the coupling advantage in the inelastic-demand subsample identifies which lever does the work. The elasticity distribution

used here, centered at negative 1.35 with 12 percent of drawn values falling into the inelastic region, reflects the range implicit in demand elasticity structures used in car rental pricing work (Li & Pang, 2017); that the coupling advantage tracks the elastic-versus-inelastic split this closely is itself evidence that price is the active ingredient.

Reported accuracy improvements for gradient-boosted and multi-task residual value models (Kim et al., 2023; Rashed et al., 2020) are measured as point-forecast error against realized resale price without reference to when disposition actually occurred relative to the model's own projected optimum, a measurement choice the seven-point RVRR gap found here calls into question: both architectures compared in this simulation share an identical projection function, so the gap traces entirely to disposition timing. A portion of the forecast error documented across that literature, including the regression-method comparisons establishing current accuracy benchmarks (Lessmann & Voß, 2017), is plausibly attributable to disposition-timing misalignment. Separating this error source from projection-function miscalibration would require rerunning those benchmark datasets against disposition dates chosen by a coupled policy, a test outside the scope of the synthetic environment used here.

Fan et al.'s (2012) dynamic programming formulation for the TxDOT fleet solves the replacement problem to exact global optimality by backward induction over a single-vehicle state space, a guarantee the alternating two-stage decomposition tested here cannot claim, because coupling the pricing subproblem into the same cycle introduces a second decision layer that backward induction alone does not resolve. The 2.4 percent mean optimality gap measured against an exact mixed-integer solver at fleet size 50 quantifies what the coupling costs in solution quality, a cost no single-vehicle exact method incurs because it never has to coordinate against a simultaneously solved pricing stage. Convergence iteration count widens from three-to-five, the band reported for the alternating solver before elasticity draws extended into the inelastic region, to three-to-six here, and the extra iteration occurs specifically when a monthly retraining event moves multiple vehicles across the elasticity branch boundary within a single daily cycle, a transition the narrower band never needed to accommodate.

CONCLUSION

These results establish, for the first time, a quantitative decomposition of fleet-management savings into a scoring-and-gating component and a price-coupling component, each measured against an identical demand and degradation environment; decoupled architectures such as Dietz and Katz's (2001) report only a single undifferentiated savings figure. That decomposition shows the price-coupling term contributing the majority share, 61 percent of the DCRD reduction observed here, a contribution no prior architecture in the reviewed literature was structured to isolate, because none combined a condition-responsive score with a reciprocal pricing-ceiling term in the same optimization cycle. What remains open is whether the collapse of the coupling advantage in the inelastic-demand regime, where price cannot serve as the coupling instrument because the ceiling structurally binds, is a general property of any price-based coupling mechanism or an artifact specific to the duration-and-distance booking filter tested here, a question resolvable only by testing an alternative, non-price coupling instrument, such as direct mileage-access control, under the identical inelastic-demand condition isolated in this simulation.

REFERENCES

- Ayuso, M., Guillén, M., & Nielsen, J. P. (2019). Improving automobile insurance ratemaking using telematics: Incorporating mileage and driver behaviour data. *Transportation*, 46(3), 735–752. <https://doi.org/10.1007/s11116-018-9890-7>
- de Albuquerque Felizola Romeral, P. A., & Zancul, E. (2025). Total cost of ownership of electric vehicles: A synthesis of critical factors. *The Journal of Engineering*, 2025, Article e70113. <https://doi.org/10.1049/tje2.70113>
- Dietz, D. C., & Katz, P. A. (2001). U S WEST implements a cogent analytical model for optimal vehicle replacement. *Interfaces*, 31(5), 65–73.
- Fan, W., Machemehl, R. B., & Gemar, M. D. (2012). Optimization of equipment replacement. *Transportation Research Record*, 2292(1), 178–184. <https://doi.org/10.3141/2292-19>
- Feng, W., & Figliozzi, M. (2014). Vehicle technologies and bus fleet replacement optimization: Problem properties and sensitivity analysis utilizing real-world data. *Public Transport*, 6(1–2), 137–157. <https://doi.org/10.1007/s12469-014-0086-z>
- Hartman, J. C., & Murphy, A. (2006). Finite-horizon equipment replacement analysis. *IIE Transactions*, 38(5), 409–419.
- Kim, M., Choi, J., Kim, J., Kim, W., Baek, Y., Bang, G., Son, K., Ryou, Y., & Kim, K.-E. (2023). Trustworthy residual vehicle value prediction for auto finance. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(13), 15537–15544. <https://doi.org/10.1609/aaai.v37i13.26842>
- Lessmann, S., Listiani, M., & Voß, S. (2010). Decision support in car leasing: A forecasting model for residual value estimation. In *Proceedings of the 31st International Conference on Information Systems (ICIS 2010)*. Association for Information Systems.
- Lessmann, S., & Voß, S. (2017). Car resale price forecasting: The impact of regression method, private information, and heterogeneity on forecast accuracy. *International Journal of Forecasting*, 33(4), 864–877. <https://doi.org/10.1016/j.ijforecast.2017.04.003>
- Li, D., & Pang, Z. (2017). Dynamic booking control for car rental revenue management: A decomposition approach. *European Journal of Operational Research*, 256(3), 850–867. <https://doi.org/10.1016/j.ejor.2016.06.048>
- Moosavi, S., & Ramnath, R. (2023). Context-aware driver risk prediction with telematics data. *Accident Analysis & Prevention*, 192, Article 107269. <https://doi.org/10.1016/j.aap.2023.107269>
- Oliveira, B. B., Carravilla, M. A., & Oliveira, J. F. (2017). Fleet and revenue management in car rental companies: A literature review and an integrated conceptual framework. *Omega*, 71, 11–26. <https://doi.org/10.1016/j.omega.2016.09.011>
- Oliveira, B. B., Carravilla, M. A., & Oliveira, J. F. (2018). Integrating pricing and capacity decisions in car rental: A matheuristic approach. *Operations Research Perspectives*, 5, 334–356. <https://doi.org/10.1016/j.orp.2018.10.002>
- Rashed, A., Jawed, S., Rehberg, J., Grabocka, J., Schmidt-Thieme, L., & Hintsches, A. (2020). A deep multi-task approach for residual value forecasting. In U. Brefeld, E. Fromont, A. Hotho, A. Knobbe, M. Maathuis, & C. Robardet (Eds.), *Machine learning and knowledge discovery in databases: ECML PKDD 2019* (Lecture Notes in Computer Science, Vol. 11908, pp. 470–485). Springer. https://doi.org/10.1007/978-3-030-46133-1_28

- Redmer, A. (2022). Strategic vehicle fleet management: A joint solution of make-or-buy, composition and replacement problems. *Journal of Quality in Maintenance Engineering*, 28(2), 327–349. <https://doi.org/10.1108/JQME-04-2020-0026>
- Waddell, R. (1983). A model for equipment replacement decisions and policies. *Interfaces*, 13(4), 1–7. <https://doi.org/10.1287/inte.13.4.1>
- Xu, M., Meng, Q., & Liu, Z. (2018). Electric vehicle fleet size and trip pricing for one-way carsharing services considering vehicle relocation and personnel assignment. *Transportation Research Part B: Methodological*, 111, 60–82. <https://doi.org/10.1016/j.trb.2018.03.001>