

SeismicShield-RL: A Preregistered and Auditable Infrastructure for Reinforcement Learning and Multi-Objective Seismic Friction-Damper Co-Design

Faramarz Kowsari

Independent Researcher, Istanbul, Türkiye

ORCID: <https://orcid.org/0000-0003-1692-0453>

Research Infrastructure / Technical Preprint

Abstract

Reinforcement learning is an attractive candidate for seismic retrofit optimization because the design space is discrete, multi-objective, computationally expensive, and affected by earthquake and structural uncertainty. The same characteristics, however, make algorithm comparisons vulnerable to unequal simulation budgets, data leakage, post-hoc model selection, source-code drift, and ambiguous restart semantics. This paper reports SeismicShield-RL, a preregistered and auditable research infrastructure for friction-damper co-design rather than a completed algorithm-ranking study. The frozen benchmark contains 136 processed Engineering Strong Motion records from 34 events, partitioned into 52 training, 20 validation, 16 pilot, and 48 confirmatory records; 16 structural states spanning 3-, 6-, 10-, and 20-story buildings; three objectives (retrofit cost proxy, maximum inter-story drift ratio, and peak floor acceleration); six stochastic methods (random search, scalar genetic algorithm, NSGA-II, PPO, IPPO, and MAPPO); and eight fixed random seeds. A runtime preflight reproduced the SHA-256 identities of all 136 processed records and successfully exercised four Tier-1 and four Tier-2 pilot fixtures without inspecting confirmatory structural-response outcomes. Deterministic planning decomposed the registered experiment into 475 atomic shards and 2,820,160 structural-response calls. Stage A alone requires 2,780,992 Tier-1 calls, corresponding to approximately 1,026.68 hours of projected sequential simulation-call time in the tested environment; the 39,168-call Tier-2 campaign projects to approximately 21.64 hours. Because the full registered computation exceeded the intended no-cost compute envelope, the full Stage-A selection campaign, the 768-item selection freeze, and confirmatory Tier-2 evaluation were deferred rather than altering the preregistered semantics. No confirmatory claim of algorithmic superiority or seismic efficacy is made. The contribution is a preserved experimental foundation that can be resumed, audited, or adapted by future researchers.

Keywords: earthquake engineering; friction dampers; reinforcement learning; multi-agent reinforcement learning; multi-objective optimization; OpenSeesPy; reproducible research; preregistration; open science

1. Introduction: Scientific Integrity Before an Algorithmic Winner

When seismic retrofit optimization is treated as a competition among algorithms, it is easy to focus on the final ranking and overlook the conditions that make the ranking defensible. A visually impressive computational result is not automatically a reproducible scientific result. In simulation-heavy research, small choices about earthquake splits, random seeds, computational budgets, checkpoint selection, retries, or source-code versions can materially affect the apparent performance of a method.

Friction dampers have a long history in passive seismic protection, where sliding friction is used to dissipate input energy [1]. More recently, reinforcement-learning methods have also been investigated for friction-damper distribution and seismic control [2]. SeismicShield-RL does not claim novelty for applying reinforcement learning, PPO, or multi-agent reinforcement learning to dampers. Its intended contribution is different: to make the comparison itself explicit, preregistered, equal-budget, auditable, and difficult to contaminate after results become known.

I began the project with a relatively simple question: can learning-based methods, particularly multi-agent approaches, discover better out-of-sample trade-offs among retrofit cost, inter-story drift, and floor acceleration than transparent search and evolutionary baselines? As the benchmark grew, a more fundamental question became

unavoidable: what would count as a fair comparison in the first place? I therefore treated reproducibility, information boundaries, execution identity, and evidence lineage as first-class research objects rather than as administrative details.

The result is unusual. The infrastructure was completed and the confirmatory boundary was preserved, but the expensive algorithm-ranking experiment was not run. This paper documents that boundary deliberately. It explains the benchmark that was frozen, what was validated, what the execution planner revealed about computational scale, why the full experiment was deferred, and how a future researcher can continue the original preregistered study without first inspecting confirmatory performance outcomes.

2. Benchmark Design

2.1 Structural ensemble

A method that performs well on a single nominal building may fail to generalize when structural properties change. To make generalization and uncertainty part of the benchmark rather than an optional demonstration, the study uses four building heights: 3, 6, 10, and 20 stories. Each height is represented by one nominal model and three frozen structural perturbations, producing 16 predefined structural states in total. The perturbations are controlled and reproducible; they are not intended to represent every source of uncertainty that can occur in a real building.

2.2 Friction-damper design space and canonical representation

At each story, an optimizer may choose a damper count from 0 through 4 and a slip-force level from the frozen grid 0, 50,000, 100,000, 200,000, or 350,000 N. A small software detail matters scientifically here. If the damper count at a story is zero, a non-zero slip-force value has no physical meaning. Without normalization, two numerical encodings could represent the same physical design while being counted as distinct candidates. The benchmark therefore canonicalizes the representation so that zero-damper stories do not acquire artificial design identities through irrelevant slip-force values.

2.3 Objectives and multi-objective interpretation

The registered problem evaluates three principal objectives: a retrofit cost proxy, maximum inter-story drift ratio (MIDR), and peak floor acceleration (PFA). MIDR is a standard engineering demand measure related to inter-story deformation, while PFA characterizes the acceleration demand transmitted to floors and is relevant to non-structural components, contents, and equipment. The benchmark does not treat either metric as a direct certification of damage or life safety.

Where an algorithm requires scalarization, the scalar objective uses prespecified weights frozen before confirmatory evaluation. The multi-objective analyses preserve Pareto information rather than converting every conclusion into a single post-hoc score. NSGA-II [6] is included specifically as a transparent evolutionary multi-objective baseline.

3. Earthquake Data and Information Boundaries

The frozen earthquake manifest contains 136 processed records from 34 events drawn from the Engineering Strong Motion database [3,4]. The scientific value of the final comparison depends not only on which records are used, but on when each partition is allowed to influence training, selection, runtime validation, or final evaluation.

Partition	Records	Registered role
Training	52	Optimizer / policy learning
Validation	20	Candidate and checkpoint selection
Pilot	16	Software and runtime validation only
Confirmatory	48	Final out-of-sample evaluation

Partition	Records	Registered role
Total	136	Frozen processed-record manifest

Training and validation data are available during Stage A. Pilot records may be used to test the software and runtime stack, but pilot outcomes are not confirmatory evidence. The confirmatory partition is reserved for the final out-of-sample structural-response evaluation. The key boundary is therefore not that confirmatory record identities can never be authenticated; the runtime preflight did verify the processed identities of all 136 records. The protected quantity is confirmatory structural-response information: confirmatory execution and performance outcomes are not authorized before the selection freeze.

This distinction is central to the design. A confirmatory comparison is interpretable only if algorithms, budgets, selected designs, learned checkpoints, and analysis rules are fixed before confirmatory outcomes are examined. The project therefore treats information flow as part of the experimental protocol.

4. Algorithms, Seeds, and Equalized Computational Budgets

The registered stochastic comparison contains six methods: random search, a scalar genetic algorithm, NSGA-II, PPO, IPPO, and MAPPO. PPO is the single-agent reinforcement-learning baseline [7]; IPPO and MAPPO represent the multi-agent comparison, with MAPPO following the centralized-training/decentralized-execution family studied by Yu et al. [8].

Eight random seeds were frozen before confirmatory evaluation: 1103, 2207, 3313, 4421, 5521, 6637, 7753, and 8861. One discipline I wanted to impose on my own workflow was to remove the option of repeatedly changing the seed set after observing disappointing runs and then reporting only a favorable realization. The seed ledger is therefore part of the scientific contract.

The same principle applies to optimization budgets. Comparing reinforcement learning with evolutionary search by iteration count would be difficult to interpret because the internal mechanics of the methods differ substantially. In this benchmark, the common budget currency is the number of structural-response calls. This does not claim that simulator calls are the only possible fairness criterion; it states the criterion chosen and frozen for this study.

5. Two-Tier Simulation Architecture

Millions of evaluations make it impractical to use the more computationally demanding structural backend for every training step. SeismicShield-RL therefore separates the research workflow into two tiers. Tier 1 is a fast research surrogate used for high-volume training and validation. Tier 2 is an OpenSeesPy-based structural-analysis backend [5] reserved for the confirmatory campaign.

The design is intentionally conservative about information flow: algorithms may train and be selected using the Stage-A pathway, but final confirmatory structural responses are evaluated only after the selection freeze. Tier 2 contains more detailed structural simulation than the fast Tier-1 path, but neither tier should be interpreted as automatically validated for a real building. The two tiers are components of a computational benchmark and remain subject to model-form and domain-specific validation limits.

6. Reproducibility Architecture

6.1 Immutable scientific source

The registered scientific implementation is identified by tag confirmatory-v0.8.2-final and exact Git commit `cecd3b6c27b5deb6cb6be7ddc478cfc407a45644`. Operational utilities may later improve packaging, transport, verification, or crash safety, but an execution presented as the original preregistered scientific study must resolve the frozen scientific implementation and contracts rather than silently substituting a new computational definition.

6.2 Cryptographic provenance for processed waveforms

File names and upstream packaging are not sufficient identifiers for a reproducible waveform pipeline. The project therefore freezes SHA-256 identities for the processed waveform objects used by the benchmark. During the preserved runtime preflight, all 136 processed-record hashes were independently reproduced and matched the frozen manifest. A hash does not certify the physical correctness of the upstream measurement; it provides a strong mechanism for detecting whether the processed numerical object used by a later execution differs from the object registered by the study.

6.3 Fail-closed execution

The confirmatory gate follows a fail-closed rule: if the expected scientific source, manifests, numerical settings, algorithm bundle, seed ledger, or analysis contract cannot be authenticated, execution should terminate explicitly rather than continue with a silent substitute. This strict behavior can be inconvenient operationally, but it makes later auditing easier because uncertainty about scientific identity is converted into a visible failure instead of an invisible branch in the computation.

7. Deterministic Execution Planning and Runtime Preflight

One of the most consequential results of the project was not an algorithmic performance metric. It was a workload calculation. Before the deterministic execution ledger was produced, I understood that the study was large, but I did not yet understand its exact computational scale. The planner transformed the registered protocol into 475 atomic shards and an exact call ledger.

Component	Shards	Structural-response calls
Tier-1 feature precompute	16	832
Tier-1 non-policy train + validate	384	1,474,560
Tier-1 learned train + validate	24	1,305,600
Stage A total	424	2,780,992
Tier-2 seeded confirmatory	48	36,864
Tier-2 support	3	2,304
Tier-2 total	51	39,168
Grand total	475	2,820,160

The runtime preflight measured a mean Tier-1 simulation-call time of approximately 1.32904380175 s/call and a mean Tier-2 time of approximately 1.98916977725 s/call in the tested environment. Applying those measurements to the frozen call counts yields approximately 1,026.68 hours of projected sequential Tier-1 simulation-call time for Stage A and approximately 21.64 hours for the Tier-2 campaign.

These numbers are not hardware-independent completion-time promises. They are projections from a measured environment. Reinforcement-learning optimization overhead, orchestration, I/O, scheduling, parallel efficiency, and hardware variation can change calendar time. The important result is that the registered Stage-A call volume, combined with measured simulation throughput, placed the study well beyond the intended no-cost compute envelope.

8. Why the Full Experiment Was Deferred

The difficult point was not the 39,168-call confirmatory campaign. It was Stage A, particularly the 24 learned-policy shards. Each learned shard contains 54,400 structural-response calls, and under the frozen implementation the training and validation behavior inside a shard is scientifically coupled. Validation occurs within the learned training loop and the selected checkpoint is part of the resulting scientific object.

I could have reduced the registered training budget. That would have made the project cheaper, but it would no longer have been the same preregistered experiment. I could also have split the learned shards simply to fit standard hosted-CI wall-clock limits. That was not a purely operational edit: it could alter the frozen execution semantics, including checkpoint-selection behavior. A small convenience run could have been useful as a demonstration, but it could also be mistaken for the confirmatory experiment if presented carelessly.

I therefore chose to stop before the expensive scientific run rather than adapt the registered experiment to the limits of free hosted compute. This is not evidence that the computation is impossible. A persistent local or self-hosted environment with sufficient resources could provide a practical execution target. The stopping decision reflects resource availability and protocol fidelity, not a failure of the scientific question.

9. What the Final Infrastructure Release Establishes

The release does not establish that MAPPO is better than PPO, IPPO, NSGA-II, scalar genetic search, or random search. Those comparisons require the registered Stage-A selection and Tier-2 confirmatory campaign, which have not been run. What the release establishes is narrower, but still useful.

- The confirmatory protocol is publicly preregistered before confirmatory structural-response evaluation.
- The scientific source, contracts, seed set, budgets, objectives, and earthquake partitions are immutably identified.
- The processed-waveform identities for all 136 records can be reproduced at the SHA-256 level.
- Four Tier-1 and four Tier-2 pilot fixtures executed successfully in the tested environment; these are runtime checks, not efficacy evidence.
- The complete registered workload has a deterministic 475-shard execution ledger totaling 2,820,160 structural-response calls.
- The selection-only workspace exposes Stage A while keeping confirmatory Tier-2 execution locked at the release boundary.
- The computational scale of the registered study is quantified by measured preflight evidence rather than by an informal estimate.

I consider the last point an engineering result in its own right. Large simulation studies should ideally measure computational feasibility before the result-producing campaign begins. Here, the preflight changed the project plan before it changed the scientific contract.

10. Reuse Value for Future Researchers

A researcher starting a similar seismic-RL study must normally make many decisions before testing a single learning algorithm: record selection, data partitions, structural task definitions, objectives, optimization budgets, seeds, checkpoint rules, provenance, solver-failure accounting, and the information boundary between selection and final testing. SeismicShield-RL does not remove the need for scientific judgment, but it provides a concrete, inspectable implementation of much of that infrastructure.

Future researchers can use the project in several ways. They can resume the original preregistered computation with more suitable computing resources; register a new protocol while retaining the benchmark structure; replace the learning algorithms while preserving the structural tasks; study the provenance and orchestration architecture independently of the seismic application; or use the repository as an example of separating exploratory engineering from confirmatory evidence. Any new study that changes the frozen protocol should be described as a new protocol rather than retrospectively relabeled as the original confirmatory experiment.

11. Scientifically Valid Continuation of the Original Study

If the original preregistered experiment is resumed, the continuation path is intentionally constrained. The purpose is to preserve the fact that confirmatory structural-response outcomes have not yet been used to tune algorithms or choose final designs.

1. Retain the immutable scientific source and frozen contracts.
2. Prepare the selection-only workspace.

3. Hydrate and use training and validation data for Stage A; do not authorize confirmatory structural-response execution or inspect confirmatory performance outcomes.
4. Execute all 424 Stage-A shards.
5. Verify the exact Stage-A total of 2,780,992 Tier-1 structural-response calls.
6. Freeze the 768 method/seed/state selections and the corresponding learned checkpoints or selected design identities.
7. Preserve hashes and provenance for selected designs and checkpoints.
8. Publish the selection freeze at an immutable reference before opening the confirmatory execution boundary.
9. Only then authorize confirmatory data hydration and Tier-2 execution.
10. Execute all 51 Tier-2 shards totaling 39,168 structural-response calls using the frozen OpenSeesPy pathway.
11. Apply the preregistered paired event-level inferential analysis and multiplicity-control rules; do not redefine the primary family after inspecting confirmatory outcomes.

The rule is simple even if the computation is not: the confirmatory response boundary should not be opened before selection is frozen and publicly attributable to an immutable scientific state.

12. Limitations

The principal limitation is explicit: the algorithm-ranking experiment remains incomplete. The current release supports reproducibility and infrastructure claims, not algorithmic efficacy claims.

Tier 1 is a fast research surrogate and must not be treated as a substitute for independent, project-specific structural analysis. Tier 2 is an OpenSeesPy-based benchmark pathway, but a future Tier-2 result would still be a computational benchmark result, not an automatic retrofit recommendation for a real building.

The 16 frozen structural states represent only the uncertainties defined by this benchmark. They do not span every source of model-form uncertainty, material variability, soil-structure interaction, construction detail, deterioration, or site condition that can matter in practice. External validity must therefore be established separately for any real engineering application.

Finally, the runtime projections are environment-dependent. Faster processors, parallel execution, different orchestration, or numerical optimization may reduce calendar time, while policy-training overhead and operational bottlenecks may increase it. The call ledger is fixed; the observed seconds per call are not universal constants.

13. Data and Code Availability

Source code, scientific contracts, documentation, execution-planning materials, release metadata, and reproducibility artifacts are publicly available. Restricted waveform bytes are not redistributed by the repository; record provenance and frozen processed-waveform hashes are used to establish the identities of the benchmark records in accordance with upstream licensing constraints.

GitHub repository: <https://github.com/FaramarzKowsari/seismicshield-rl>

Project website: <https://faramarzkowsari.github.io/seismicshield-rl/>

Researcher Guide: <https://faramarzkowsari.github.io/seismicshield-rl/researcher-guide.html>

Zenodo exact software version DOI: <https://doi.org/10.5281/zenodo.22067278>

Zenodo Concept DOI (software record lineage): <https://doi.org/10.5281/zenodo.22067277>

OSF preregistration DOI: <https://doi.org/10.17605/OSF.IO/64DTX>

14. Engineering and Safety Scope

SeismicShield-RL is research software. It does not certify buildings, prescribe construction, demonstrate regulatory or building-code compliance, establish life-safety performance, or replace project-specific analysis and review by appropriately qualified and, where required, licensed structural engineers. Any future application of ideas from this benchmark to a real structure would require independent structural modeling, verification, engineering review, and compliance with the applicable standards and legal responsibilities.

15. AI-Assistance Disclosure

Generative-AI tools were used during parts of the project for software-development assistance, document organization, and preparation of supporting materials. For this manuscript, the author drafted and substantively revised the scientific narrative in Persian; generative-AI tools were then used for English translation, copy-editing, terminology consistency, and document formatting. The author reviewed the scientific content and accepts responsibility for the claims, interpretations, and final manuscript. Numerical claims in this paper are tied to frozen project artifacts, and no confirmatory structural-response outcome was generated or inspected for algorithm selection or final efficacy reporting.

16. Conclusion

I began SeismicShield-RL by asking whether modern reinforcement-learning methods could outperform more conventional optimization approaches for friction-damper co-design. The project eventually forced a second question to come first: could I define and preserve a comparison that I would still consider defensible after the results were known?

That change in priority shaped the final release. The earthquake partitions, structural tasks, objectives, algorithms, seeds, budgets, scientific source, provenance checks, deterministic execution plan, and confirmatory information boundary were frozen before the final algorithm-ranking experiment. Runtime preflight then made the computational cost concrete. Rather than reduce the registered budget or modify the frozen semantics to fit the available no-cost infrastructure, I stopped before Stage A and the confirmatory campaign.

The project therefore ends without an algorithmic winner. It also ends without contaminating the confirmatory structural-response boundary through outcome inspection. That distinction is the main reason the infrastructure remains scientifically useful: a future researcher with suitable resources can still answer the original question without first having tuned the study to the final outcomes.

SeismicShield-RL should consequently be read as a preserved experimental foundation, not as evidence that any particular learning or optimization method is superior. The final empirical question remains open. The purpose of this release is to make that future answer easier to audit, reproduce, and interpret.

References

- [1] Pall, A. S., & Marsh, C. (1982). Response of friction damped braced frames. *Journal of the Structural Division, ASCE*, 108(9), 1313-1323. <https://doi.org/10.1061/JSDEAG.0005968>
- [2] Kurucu, M. C., Atam, E., Guzelkaya, M., & Eksin, I. (2024). Intelligent Computational Methods for Optimal Distribution of Friction Dampers in Seismic Protection of Buildings. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 8(4), 3055-3066. <https://doi.org/10.1109/TETCI.2024.3369909>
- [3] Luzi, L., Lanzano, G., Felicetta, C., D'Amico, M. C., Russo, E., Sgobba, S., Pacor, F., & ORFEUS Working Group 5. (2020). Engineering Strong Motion Database (ESM), Version 2.0. Istituto Nazionale di Geofisica e Vulcanologia. <https://doi.org/10.13127/ESM.2>
- [4] Lanzano, G., Sgobba, S., Luzi, L., Puglia, R., Pacor, F., Felicetta, C., D'Amico, M., Cotton, F., & Bindi, D. (2019). The pan-European Engineering Strong Motion (ESM) flatfile: compilation criteria and data statistics. *Bulletin of Earthquake Engineering*, 17, 561-582. <https://doi.org/10.1007/s10518-018-0480-z>
- [5] Zhu, M., McKenna, F., & Scott, M. H. (2018). OpenSeesPy: Python library for the OpenSees finite element framework. *SoftwareX*, 7, 6-11. <https://doi.org/10.1016/j.softx.2017.10.009>
- [6] Deb, K., Pratap, A., Agarwal, S., & Meyarivan, T. (2002). A fast and elitist multiobjective genetic algorithm: NSGA-II. *IEEE Transactions on Evolutionary Computation*, 6(2), 182-197. <https://doi.org/10.1109/4235.996017>
- [7] Schulman, J., Wolski, F., Dhariwal, P., Radford, A., & Klimov, O. (2017). Proximal Policy Optimization Algorithms. *arXiv:1707.06347*. <https://doi.org/10.48550/arXiv.1707.06347>
- [8] Yu, C., Velu, A., Vinitsky, E., Gao, J., Wang, Y., Bayen, A., & Wu, Y. (2022). The Surprising Effectiveness of PPO in Cooperative Multi-Agent Games. *Advances in Neural Information Processing Systems*, 35, 24611-24624. <https://doi.org/10.52202/068431-1787>
- [9] Kowsari, F. (2026). SeismicShield-RL v0.8.2 - Final Infrastructure Release [Software]. Zenodo. <https://doi.org/10.5281/zenodo.22067278>
- [10] Kowsari, F. (2026). SeismicShield-RL preregistered protocol. OSF. <https://doi.org/10.17605/OSF.IO/64DTX>