

Peer Review: A cloud-native, compute-optimized tile grid for land cover analysis

Thomas Maschler¹ , Ryan C. McCarthy¹ , & Christopher B. Anderson¹ 

Collaborators: 3 reviewers

¹ Planet Labs PBC, San Francisco, CA, USA

Transparent Peer Review

- [View reviewer summaries](#)
- [View initial submission with reviewer comments](#)

Recommended Citation

Maschler, T., R. C. McCarthy, and C. B. Anderson. 2026. A cloud-native, compute-optimized tile grid for land cover analysis. Stacks Journal: 26010. <https://doi.org/10.60102/stacks-26010>.

Accepted by 2 of 3 reviewers

Conflicts of Interest

All authors are employees of Planet Labs PBC, which funded this work.

Publishing History

Submitted May 5 2026

Accepted July 7 2026

Published September 4 2026

Corresponding Author

Christopher B. Anderson
christopher@planet.com



Open Access



Peer-Reviewed



Creative Commons



Reviewer Summaries

Tim Bailey

Initial Submission

Do you have any conflicts of interest that could bias your ability to provide an independent review?

no

What did the authors do a good job with?

Mostly. I really appreciated the technical approach. However the authors should have more in depth discussion of the OGC DGGs standards which are emerging to solve similar problems.

How do you think this research will contribute to the field?

The technical implementation is interesting. I could see conducting benchmark tests on this approach, however it is not clear that the design will convince large data operations to adopt these methods based on the material in the paper.

Regarding the study design and methods, what do the authors need to fix or improve upon to be fit for publication?

I would specifically compare this method to more discrete global grids besides H3. The tradeoffs present in projection selection have been challenging geographers for centuries. The paper could benefit from better engagement with DGGs systems that are built for similar user stories.

Regarding the analysis and interpretation of their findings, what do the authors need to fix or improve upon to be fit for publication?

The technical decisions are reasonable for selecting projection methods for a project but the case for wide adoption was not convincing.

Is there anything else you think the authors need to fix in their article to be fit for publication?

No answer provided

Do you have any concerns about the ethics of this research?

No

Based on your review, what should happen next?

This paper requires minor revisions but does not need further peer review

Would you like to be listed as a Collaborator on the final publication?

No, I do not want to be listed as a Collaborator

John Hogland

Initial Submission

Do you have any conflicts of interest that could bias your ability to provide an independent review?

No

What did the authors do a good job with?

Describing the gridded tiling schem.

How do you think this research will contribute to the field?



Provide an additional projection that can be used to better address area and shape for substantial portions of the globe.

Regarding the study design and methods, what do the authors need to fix or improve upon to be fit for publication?

The manuscript would benefit from a more formal analysis. See comments in the review.

Regarding the analysis and interpretation of their findings, what do the authors need to fix or improve upon to be fit for publication?

While the authors stipulate that, "the choice of grid structure, projection, and indexing scheme has significant downstream implications ...", the choice of projection depends on the analyses being performed. See comments in the manuscript.

Is there anything else you think the authors need to fix in their article to be fit for publication?

I believe the emphasis on this manuscript should be on describing the benefits of the discussed projection. The comparison in this case should focus more on pros and cons of the described projection with less of an emphasis on if a given projection is better or worse than the other projection.

Do you have any concerns about the ethics of this research?

No

Based on your review, what should happen next?

This paper requires minor revisions but does not need further peer review

Would you like to be listed as a Collaborator on the final publication?

No, I do not want to be listed as a Collaborator

Pranath Reddy Kumbaum

Initial Submission

Do you have any conflicts of interest that could bias your ability to provide an independent review?

no

What did the authors do a good job with?

The paper addresses a real and well-motivated problem in Earth observation data processing, namely the fragmentation, resolution mismatch, and repeated resampling that arise from inconsistent global tiling schemes, and it situates this clearly against prior work including DGGs, Equi7Grid, Tile Matrix Sets, and EASE-Grid. Its core technical observation is sensible and mildly novel: a base-3 nonary hierarchy approximates common EO sensor resolutions (1, 3, 10, 30 m) more closely than the factor-of-two quad-tree schemes dominant in web mapping. The work delivers a concrete, reproducible artifact, a Tile Matrix Set specification with working pyproj and morecantile code, which is a genuine adoption and reproducibility strength. The prefix-based tile index is practical, encoding hierarchy and ancestry directly, supporting prefix-scan metadata lookups without spatial queries, and remaining interoperable with standard (z, x, y) addressing. Finally, the paper is candid about its tradeoffs, repeatedly framing the grid as workload-specific rather than universal and listing its limitations explicitly.

How do you think this research will contribute to the field?

This research is most likely to contribute as practical infrastructure rather than as a scientific or ML advance. Its main value is a reproducible, standards-compatible tiling convention, namely a concrete EASE-Grid Tile Matrix Set with a base-3 nonary hierarchy and a prefix-based index, packaged with working



code, which teams building large-scale area-based EO pipelines can adopt to reduce resampling, stabilize compute and storage costs, and simplify metadata management and spatial train/test splitting. The base-3 alignment with common sensor resolutions is a useful design insight that others may borrow even without adopting this exact grid, and the documented design framework (evaluating projection footprint, distortion, layout, and tooling compatibility against explicit operational requirements) offers a reusable template for groups designing grids for their own workflows. Its contribution to machine learning specifically is more limited, since the ML benefits are argued rather than demonstrated, so the durable impact is best understood as adoption-and-template value within the applied EO community rather than a shift in ML understanding or capability.

Regarding the study design and methods, what do the authors need to fix or improve upon to be fit for publication?

Regarding study design and methods, the most important fix is the gap between framing and evidence: the work is presented as "machine learning-optimized" yet includes no ML component (no model, no downstream task, no comparison of learning outcomes across grids), so the authors should either add at least one concrete experiment, for example training an identical model on data tiled by a quad-tree baseline versus the proposed nonary EASE-Grid and reporting both accuracy and resampling volume alongside a spatial-leakage comparison, or substantially soften the ML claims and retile the work as area-based EO infrastructure. The projection-evaluation methods also need tightening: the distortion analysis uses an unconventional pentagram variant of Tissot's indicatrix with coefficient-of-variation aggregation that is never justified and should be explained or replaced with a conventional metric and reported in tabular form, while the qualitative Table 2 scoring relies on an unweighted emoji rubric with no rater protocol and should be replaced with explicit weighted criteria so the final selection is traceable, and the selection rationale should be made honest that EASE-Grid was chosen over the geometrically comparable Gall-Peters largely on registry availability rather than geometry.

Regarding the analysis and interpretation of their findings, what do the authors need to fix or improve upon to be fit for publication?

Regarding analysis and interpretation, the authors overread results that are largely deterministic geometry as if they were validated performance findings, and this should be corrected throughout. The tile-count comparison (Table 1) mostly re-derives well-known projection facts, such as Web Mercator requiring roughly twice the pixels due to high-latitude area inflation, yet tile count is presented as a proxy for compute and storage cost without accounting for compression, boundary-tile clipping and sparsity, or actual IO patterns, so the cost interpretation should be hedged or substantiated with real storage and runtime measurements. The interpretation of distortion is similarly strained: the authors conclude EASE-Grid is the most balanced choice while their own analysis shows it maximizes shape and angular distortion at high latitudes, and they wave this away as "predictable" without engaging the consequence they themselves raise, that this distortion is exactly what degrades vision-model generalization, so the interpretation needs to confront rather than minimize that tradeoff and scope the "optimal" claim explicitly to area-based workflows. The most ML-relevant interpretive claims, that grid-aligned sampling reduces train/test leakage and that sibling cells provide useful change-detection context, are stated as established outcomes in the "emergent use cases" section but rest on no measurement, and these should be reframed as hypotheses or supported with even a minimal quantitative result. Finally, the conclusion that the realized resolutions are "negligibly" different from nominal EO scales is an interpretation contradicted by the authors' own numbers (a 15 to 27 percent gap at the fine end), so this should be quantified honestly, with the interpretation acknowledging that resampling is reduced rather than



eliminated and reporting the residual against a quad-tree baseline rather than asserting practical equivalence.

Is there anything else you think the authors need to fix in their article to be fit for publication?

No

Do you have any concerns about the ethics of this research?

No

Based on your review, what should happen next?

This paper needs major revisions and another round of peer review