

Cross-Modal Generative Models for Remote Sensing Image Enhancement and Earth Observation Applications

Roy A. Freeman

School of Information Technology, University of Cincinnati, Cincinnati, OH, USA.
roy.freeman@uc.edu

Anirudh Jha

Department of Electrical Engineering and Computer Science, University of Kansas, Lawrence,
KS, USA.
anirudhjha366@ku.edu

Meateo Woods

School of Electrical Engineering and Computer Science, Oregon State University, Corvallis,
OR, USA.
contactmatteo@oregonstate.edu

Bastian Alvarez

Department of Computer Science and Engineering, University of Nevada, Reno, Reno, NV,
USA.
bastian.alvarez@unr.edu

Abstract

Cross-modal generative models are becoming a central architectural paradigm for remote sensing image enhancement and large-scale Earth observation workflows. These models integrate heterogeneous data sources, including multispectral imagery, synthetic aperture radar, text-based environmental reports, and geospatial metadata, to reconstruct degraded scenes, fill missing observations, and generate physically plausible interpretations of the Earth surface. This paper provides a systems-level analysis of cross-modal generative approaches in remote sensing, focusing not on narrow algorithmic performance but on structural trade-offs, data governance, deployment architectures, robustness, fairness, sustainability, and policy implications. It examines how adversarial and diffusion-based generative mechanisms are reorganized within Earth observation pipelines and how they interact with transformer-based representation learning. The discussion emphasizes the importance of multi-sensor data governance, calibration under distribution shift, edge-to-cloud deployment, and the regulatory pressures emerging around foundation models. The paper further considers how generative outputs are increasingly used in humanitarian mapping, climate monitoring, and environmental policy, requiring careful attention to uncertainty communication and equitable access. By connecting architectural decisions to broader infrastructural and societal constraints, the paper argues that the next generation of Earth observation systems must treat generative enhancement not merely as an image reconstruction task but as a governed socio-technical infrastructure.

Keywords

cross-modal generative models, remote sensing, Earth observation, image enhancement, deployment architectures, data governance, foundation models, fairness, sustainability.

1. Introduction

Remote sensing systems now produce large volumes of heterogeneous observations across optical, thermal, radar, and topographical modalities. The integration of deep learning into these pipelines has reshaped the analysis of land cover, vegetation dynamics, urban expansion, and disaster response at regional and planetary scales. Early reviews of deep learning in remote sensing established that convolutional and recurrent architectures could outperform traditional feature engineering for classification, segmentation, and change detection when sufficient labeled data were available [1]. Subsequent meta-analyses confirmed that deep learning methods have expanded across nearly every major remote sensing application area, while also revealing persistent bottlenecks related to label scarcity, domain shift, and the interpretability of learned representations [2]. Within this context, generative models have moved from an experimental curiosity to a practical instrument for image enhancement, super-resolution, cloud removal, and synthetic data generation.

The emergence of cross-modal generative models introduces a deeper transformation. Instead of operating only within a single image modality, these models learn shared latent spaces that connect electromagnetic observations with textual descriptions, temporal sequences, and geospatial context. This cross-modal capacity enables remote sensing systems to reconstruct missing observations using cues from auxiliary data sources and to produce outputs that are not merely sharpened versions of the input but semantically conditioned reconstructions. The structural implications of this shift are significant. Earth observation pipelines must now be designed around multimodal encoders, generative decoders, alignment modules, and uncertainty mechanisms that extend well beyond classical image processing. At the same time, the deployment of such models raises new governance questions about data provenance, model bias, energy consumption, calibration, and the responsible use of generated environmental imagery. This paper addresses these system-level concerns by examining the architectural foundations, infrastructure requirements, deployment alternatives, robustness constraints, fairness considerations, and policy frameworks that shape cross-modal generative model use in Earth observation.

2. Cross-Modal Generative Foundations for Remote Sensing

Generative adversarial networks established the core principle of learning a generator that produces realistic samples by competing with a discriminator [3]. In remote sensing, conditional variants of this framework allowed image translation tasks such as super-resolution, pan-sharpening, and synthetic aperture radar to optical image reconstruction to be formulated as mappings from an observed input domain to a target domain [4]. The conditional adversarial objective provides a flexible structural mechanism for preserving high-frequency spatial detail while constraining outputs according to paired or unpaired translation settings. Unpaired image translation architectures further relaxed the data requirements by using cycle consistency to align domains without requiring pixel-level paired examples [5]. This capability is particularly relevant for remote sensing because simultaneous acquisitions across sensors are often temporally misaligned or affected by different atmospheric conditions.

Diffusion-based generative models have recently extended this capacity by decomposing image generation into a sequence of denoising transformations. Latent diffusion architectures

perform the generative process in a compressed latent space, reducing computational demands while preserving reconstruction quality and enabling conditioning on multiple modalities [6]. Denoising diffusion probabilistic models formalize the iterative refinement mechanism that has become a standard alternative to adversarial training for high-resolution synthesis [7]. For Earth observation, diffusion models offer structural advantages when reconstructing missing regions or removing atmospheric distortions because the probabilistic refinement process can be conditioned on auxiliary variables such as cloud masks, seasonal indicators, or terrain elevation. The generative process can therefore be arranged as a controlled restoration problem rather than an unconditional synthesis task.

The cross-modal dimension emerges when these generative components are coupled with representation learning systems that align visual and textual or temporal information. Generator outputs can be conditioned on language descriptions of land use, weather events, or humanitarian priorities, enabling analysts to query generative enhancement workflows through semantically meaningful interfaces. This alignment also introduces structural complexity because the system must balance the fidelity of the generated image with the semantic consistency required by the conditioning signal. In operational terms, this shifts the design question from selecting a generative architecture to designing a joint embedding and generation infrastructure in which sensor data, textual metadata, and spatiotemporal context are continuously co-aligned.

3. Architectures and System-Level Design Trade-offs

The dominant computational building blocks of contemporary cross-modal systems derive from attention mechanisms that allow models to relate distant spatial and temporal positions without relying on fixed receptive fields [8]. Vision transformers adapted this principle to image data by treating images as sequences of patches, enabling more flexible spatial reasoning than convolutional baselines when sufficient data are available [9]. Cross-modal foundation models built on contrastive language-image training demonstrated that a single visual encoder can be aligned with a broad textual concept space, producing representations that transfer across many downstream tasks [10]. In remote sensing, these architectural shifts create a tension between generic visual representations learned from natural imagery and the specialized spectral, temporal, and geometric characteristics of Earth observation data.

The remote sensing community has responded by developing foundation models specifically pretrained on multispectral and temporal satellite imagery. Masked autoencoding strategies for satellite image sequences exploit temporal and spectral redundancy to learn representations that can be fine-tuned for crop monitoring, segmentation, and change detection [11]. RingMo introduced masked image modeling on remote sensing scenes to improve representation quality across varied geographic and sensor contexts [12]. SatlasPretrain further contributed large-scale pretraining data across multiple modalities and tasks, reinforcing the importance of broad data coverage for downstream generalization [13]. These specialized pretraining efforts demonstrate a critical system-level decision: generic cross-modal models can leverage language supervision and internet-scale knowledge, but they may require substantial adaptation to account for the physical properties of remote sensing instruments.

The trade-off between generic and domain-specialized architectures is not purely technical. Generic cross-modal models may accelerate prototyping and support interactive Earth observation tools because analysts can express requests in natural language. However, their outputs can also introduce spectral inaccuracies or geographic inconsistencies when used for

scientific analysis. Domain-specialized remote sensing models may be more physically consistent but often lack the breadth of semantic alignment needed for interactive cross-modal queries. A robust Earth observation infrastructure therefore requires a modular architecture in which specialized remote sensing encoders are paired with cross-modal alignment heads, generative decoders, and calibration components. Such an architecture permits distinct failure modes to be isolated and governed rather than embedded invisibly within a single monolithic model.

4. Data Infrastructure, Multi-Sensor Fusion, and Training Governance

Cross-modal generative enhancement depends on data infrastructure that can harmonize observations from satellites, airborne platforms, ground sensors, and textual environmental records. The value of satellite data for socioeconomic measurement has been demonstrated in poverty prediction frameworks that combine daytime and nighttime imagery with survey data, indicating that remote sensing models can support policy-relevant inference beyond purely physical mapping [14]. Such applications place additional demands on data governance because the model must preserve geographic variation while avoiding overgeneralization from training regions to culturally and economically diverse areas.

Multi-sensor fusion also raises questions about the provenance and quality of training examples. Generative models are sensitive to missing data mechanisms, calibration differences across instruments, and inconsistencies in georeferencing. Data infrastructure must therefore include metadata standards that document sensor type, acquisition time, cloud cover, preprocessing history, and licensing constraints. Without such governance, a generative model trained on heterogeneous sources may learn spurious correlations that undermine scientific validity. Earth system science has emphasized the need for deep learning methods to be integrated with process understanding rather than used as black-box predictors, highlighting the importance of domain-aware training and evaluation [15]. Interdisciplinary research agendas similarly call for collective governance of Earth observation data and models so that machine learning outputs can be trusted in environmental decision contexts [16].

These governance considerations are intensified by the cross-modal nature of modern systems. Textual prompts, auxiliary geospatial datasets, and user feedback loops all become part of the training and inference ecosystem. When a model learns to associate textual descriptions with image features, errors in language understanding can propagate into spatial reconstructions. Training governance must therefore cover not only image data quality but also the curation of paired text, the representation of geographic concepts, and the procedures for updating models as new sensors or labels become available. A system that treats these dimensions as separate will likely produce enhancements that are visually compelling but unreliable for operational Earth observation.

5. Deployment Architectures and Edge-to-Cloud Integration

Generative enhancement models are computationally intensive, yet many remote sensing applications require low-latency outputs near the point of acquisition or in constrained field environments. The shift toward edge-native generative foundation models reflects an emerging response to this challenge, with recent technical reports describing text-to-image systems trained entirely on China-developed AI accelerators and optimized for edge execution [17]. Although this cited work focuses on text-to-image generation rather than remote sensing specifically, its architecture and training strategy are directly relevant to Earth

observation because they illustrate the feasibility of moving generative workloads away from centralized GPU clusters and toward alternative hardware ecosystems. For remote sensing, edge deployment may support onboard cloud removal, rapid disaster mapping, or privacy-preserving processing of sensitive geographic imagery.

The edge-to-cloud integration problem involves more than model compression. It requires decisions about which components of the generative pipeline should reside on the sensor or edge node, which should be executed in regional cloud facilities, and how intermediate representations should be transmitted. A common system design is to place lightweight encoders and task-specific decoders at the edge while retaining large cross-modal foundation models in the cloud. Alternatively, federated training or decentralized fine-tuning can be used to update models using local observations without transferring raw imagery. These architectural choices have direct implications for communication bandwidth, energy consumption, and resilience during connectivity interruptions. Earth observation systems operating in disaster zones, maritime environments, or remote agricultural regions must therefore be designed with graceful degradation: the edge component should still produce useful outputs even when cloud-based cross-modal conditioning is unavailable.

Deployment architectures also interact with technological sovereignty and supply chain resilience. The reliance on specific accelerator families, cloud providers, or model formats can shape national and institutional capacity to maintain critical Earth observation infrastructure. The development of edge-native models on alternative hardware demonstrates that generative capabilities are not inherently bound to a single vendor ecosystem [17]. However, achieving comparable performance across diverse accelerators requires substantial investment in compiler toolchains, precision reduction methods, and deployment frameworks. System designers must weigh the benefits of specialized hardware against the portability and long-term maintainability of models that may need to be redeployed across multiple agencies and jurisdictions.

6. Robustness, Calibration, and Sustainability

The use of generative models for remote sensing enhancement introduces distinctive robustness challenges. A model trained to reconstruct missing optical pixels may generate plausible structures that are not supported by the underlying physical measurement. In scientific workflows, this can lead to confirmation bias if analysts treat enhanced imagery as raw observation. Robustness therefore requires explicit calibration mechanisms that report uncertainty, distinguish interpolated from observed regions, and flag outputs that fall outside the training distribution. Earth observation systems must also monitor distribution shift caused by seasonal changes, land cover transitions, sensor degradation, and climate variability. A generative enhancement model that performs well on temperate agricultural scenes may fail in arid or polar environments unless the training infrastructure explicitly accounts for geographic diversity.

Sustainability has become an equally important constraint. Training large cross-modal models consumes substantial computational energy, and repeated inference for global-scale remote sensing products can amplify these costs. Quantitative analyses of machine learning carbon emissions have shown that model architecture, data center location, and hardware efficiency significantly affect the total environmental footprint [18]. Energy policy research in deep learning further demonstrates that efficiency gains alone may not reduce aggregate energy use if model scale and application scope continue to expand [19]. For Earth observation, this creates a structural tension: generative enhancement can improve the usability of existing

satellite data and reduce the need for repeated acquisitions, but it also consumes energy to produce augmented imagery. Life cycle evaluation should therefore consider both the operational savings from better data products and the environmental cost of training and deployment.

A sustainability-oriented system design may prioritize modular fine-tuning rather than full retraining, use low-rank adaptation, or employ distillation to reduce inference costs. It may also schedule computationally intensive generative tasks during periods of renewable energy availability or restrict high-resolution enhancement to regions where decisions require such detail. These choices are not merely operational details; they determine whether cross-modal generative models can be justified as components of long-term environmental monitoring systems.

7. Data Documentation, Fairness, and Policy Implications

The adoption of cross-modal generative models in public sector Earth observation raises urgent questions about data documentation and accountability. Datasheets for datasets have been proposed as a mechanism for ensuring that dataset creators document the motivation, composition, collection process, and recommended uses of training data [20]. In remote sensing, such documentation is essential because imagery may be sourced from commercial satellites, government archives, or humanitarian campaigns, each with distinct licensing, privacy, and representational constraints. A generative model trained on uneven geographic coverage may produce high-quality enhancements for well-documented regions while failing in underrepresented areas. Datasheet-style governance can help operators identify these gaps before deployment and communicate limitations to downstream users.

Foundation models have introduced broader concerns because their pretraining data and capabilities are often opaque, and their failures can propagate across many downstream applications [21]. When cross-modal generative systems are used for environmental policy, spatial planning, or disaster response, the consequences of biased outputs can be distributed unevenly across communities. Fairness considerations require more than aggregate accuracy metrics; they demand analysis of model performance across land tenure regimes, urban and rural contexts, and regions with different data availability. Fair machine learning research has emphasized that technical definitions of fairness may be insufficient when institutional and social context is ignored [22]. Earth observation systems must therefore embed fairness assessment within their evaluation pipelines, using participatory methods and independent audits rather than relying solely on automated metrics.

Policy frameworks are beginning to address these issues through transparency requirements, impact assessments, and procurement standards for public sector artificial intelligence. Remote sensing agencies and research institutions that deploy generative enhancement models will need to document how generated imagery is labeled, how uncertainty is communicated, and how users can contest automated outputs. The alignment between technical architectures and policy requirements is not automatic. It requires organizational capacity, interdisciplinary review, and legal clarity about the distinction between observed data and model-generated reconstructions.

8. Future Directions and Emerging Infrastructures

Future cross-modal Earth observation systems will likely integrate multiple generative modalities, including text-conditioned image enhancement, radar-to-optical translation, temporal gap filling, and three-dimensional surface reconstruction. These capabilities will

converge around shared geospatial foundation models that are actively aligned with physical sensor models and environmental knowledge graphs. The generative outputs of such systems may be used not only for visual analysis but also as inputs to hydrological models, agricultural forecasts, and climate adaptation planning. This increases the need for traceability: every generated pixel must be associated with its conditioning inputs, model version, calibration status, and governance record.

Emerging infrastructures may also support decentralized model updates, where regional agencies fine-tune shared foundation models using local observations and periodically contribute model deltas to a federated repository. Such arrangements could improve geographic equity and reduce data transfer costs, but they also require new protocols for versioning, quality control, and accountability. The evolution of edge-native generative training further suggests that future Earth observation platforms may combine onboard enhancement with selective cloud-based semantic conditioning, reducing raw data transmission and improving responsiveness.

The long-term success of these systems depends on their institutional integration. Cross-modal generative models should be governed through data consortia, standards bodies, and public oversight mechanisms that connect satellite operators, scientific users, humanitarian organizations, and affected communities. Without this integration, the technical capacity of generative models may outpace the governance structures needed to ensure their responsible use. Earth observation is both a scientific and a public good, and its infrastructures must reflect that duality.

9. Conclusion

Cross-modal generative models are reshaping the architecture of remote sensing image enhancement and Earth observation. Their ability to integrate spectral, temporal, textual, and geospatial cues creates new opportunities for reconstructing degraded observations, filling data gaps, and enabling interactive environmental analysis. However, these capabilities introduce structural challenges that cannot be resolved by model optimization alone. The design of Earth observation systems must account for architectural trade-offs between generic and domain-specialized models, multi-sensor data governance, edge-to-cloud deployment, robustness under distribution shift, energy sustainability, fairness, and public accountability. The growing use of specialized remote sensing foundation models and edge-native generative platforms indicates that the field is moving toward more distributed and multimodal infrastructures. For these infrastructures to be trustworthy, they must be embedded within governance frameworks that document data provenance, communicate uncertainty, and support equitable access. The future of Earth observation will therefore depend as much on institutional design and policy development as on advances in generative modeling.

References

1. Zhu, X. X., Tuia, D., Mou, L., Xia, G.-S., Zhang, L., Xu, F., & Fraundorfer, F. (2017). Deep learning in remote sensing: A comprehensive review and list of resources. *IEEE Geoscience and Remote Sensing Magazine*, 5(4), 8–36.
2. Ma, L., Liu, Y., Zhang, X., Ye, Y., Yin, G., & Johnson, B. A. (2019). Deep learning in remote sensing applications: A meta-analysis and review. *ISPRS Journal of Photogrammetry and Remote Sensing*, 152, 166–177.

3. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014). Generative adversarial nets. *Advances in Neural Information Processing Systems*, 27, 2672–2680.
4. Isola, P., Zhu, J.-Y., Zhou, T., & Efros, A. A. (2017). Image-to-image translation with conditional adversarial networks. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 1125–1134.
5. Zhu, J.-Y., Park, T., Isola, P., & Efros, A. A. (2017). Unpaired image-to-image translation using cycle-consistent adversarial networks. *Proceedings of the IEEE International Conference on Computer Vision*, 2223–2232.
6. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., & Ommer, B. (2022). High-resolution image synthesis with latent diffusion models. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 10684–10695.
7. Ho, J., Jain, A., & Abbeel, P. (2020). Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems*, 33, 6840–6851.
8. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
9. Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., & Houlsby, N. (2021). An image is worth 16x16 words: Transformers for image recognition at scale. *International Conference on Learning Representations*.
10. Radford, A., Kim, J. W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., & Sutskever, I. (2021). Learning transferable visual models from natural language supervision. *Proceedings of the 38th International Conference on Machine Learning*, 8748–8763.
11. Cong, Y., Khanna, S., Meng, C., Liu, P., Rozi, E., He, Y., Burke, M., Lobell, D., & Ermon, S. (2022). SatMAE: Pre-training transformers for temporal and multi-spectral satellite imagery. *Advances in Neural Information Processing Systems*, 35, 197–211.
12. Sun, X., Wang, P., Lu, W., Zhu, Z., Lu, X., He, Q., Li, J., Rong, X., Yang, Z., Chang, H., & Fu, K. (2022). RingMo: A remote sensing foundation model with masked image modeling. *IEEE Transactions on Geoscience and Remote Sensing*, 60, 1–22.
13. Bastani, F., Wolters, P., Gupta, R., Ferdinando, J., & Kembhavi, A. (2023). SatlasPretrain: A large-scale dataset for remote sensing image understanding. *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 16772–16782.
14. Jean, N., Burke, M., Xie, M., Davis, W. M., Lobell, D. B., & Ermon, S. (2016). Combining satellite imagery and machine learning to predict poverty. *Science*, 353(6301), 790–794.
15. Reichstein, M., Camps-Valls, G., Stevens, B., Jung, M., Denzler, J., Carvalhais, N., & Prabhat. (2019). Deep learning and process understanding for data-driven Earth system science. *Nature*, 566(7743), 195–204.
16. Tuia, D., Roscher, R., Wegner, J. D., Jacobs, N., Zhu, X. X., & Camps-Valls, G. (2021). Toward a collective agenda on AI for Earth science data analysis. *IEEE Geoscience and Remote Sensing Magazine*, 9(2), 88–104.

17. Chen, Ce, et al. "JuZhou 1.0 Technical Report: The First Edge-Native Text-to-Image Foundation Model Trained Entirely on China-Developed AI Accelerators." arXiv preprint arXiv:2606.28421 (2026).
18. Lacoste, A., Luccioni, A., Schmidt, V., & Dandres, T. (2019). Quantifying the carbon emissions of machine learning. arXiv preprint arXiv:1910.09700.
19. Strubell, E., Ganesh, A., & McCallum, A. (2019). Energy and policy considerations for deep learning in NLP. Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics, 3645–3650.
20. Gebru, T., Morgenstern, J., Vecchione, B., Vaughan, J. W., Wallach, H., Daumé III, H., & Crawford, K. (2021). Datasheets for datasets. Communications of the ACM, 64(12), 86–92.
21. Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., & Liang, P. (2021). On the opportunities and risks of foundation models. arXiv preprint arXiv:2108.07258.
22. Veale, M., & Binns, R. (2017). Fairer machine learning in the real world: Mitigating discrimination without collecting disparate impact. Big Data & Society, 4(2), 2053951717743530.