

Scaling seismic Interpretation with Foundation Models

Henri Houllevigue^{1*}, Susana Tierrablanca¹, Yermek Balabekov¹, Olga Brusova¹, Vi Ly¹, Meher Gajula¹, Ben Lasscock¹ and Alejandro Valenciano¹ examine how the same paradigm can support broader geological interpretation and discuss representative applications, including structural interpretation, depositional-element and geobody detection, and seismic facies classification.

Abstract

Seismic interpretation is a multiscale geological workflow that integrates structure, stratigraphy, depositional patterns, seismic texture, well control and uncertainty. Task-specific machine-learning methods have accelerated individual steps, but they often require labelled data and retraining for each new application or dataset. Seismic foundation models offer a more scalable alternative by learning reusable 3D seismic representations from large volumes of unlabeled data.

Building on prior work demonstrating scalable seismic foundation-model pretraining and salt segmentation, this article examines how the same paradigm can support broader geological interpretation. We discuss representative applications, including structural interpretation, depositional-element and geobody detection, and seismic facies classification. Across these examples, model outputs are best understood as probability volumes and interpretation candidates rather than final geological products.

We argue that foundation models shift AI from isolated task automation toward foundation-assisted interpretation. Their practical value lies not only in faster prediction but also in broader screening, improved repeatability, and more effective use of interpreter time. Interpreters remain essential for quality control, geological calibration, uncertainty assessment, and final interpretation.

Introduction

Seismic interpretation has always been a problem of pattern recognition under uncertainty. The interpreter must connect reflections across space, determine whether discontinuities are faults or processing artifacts, distinguish stratigraphic terminations from noise, and infer geological meaning from band-limited seismic images. The tools have changed dramatically, from paper sections and coloured pencils to 3D workstations, seismic attributes, quantitative interpretation and machine learning, although the central challenge remains the same: how to turn a seismic image into a coherent geological model.

Modern seismic datasets have made this challenge more acute. A single 3D survey may contain many terabytes of

stacked and pre-stack data, multiple processing versions, and dozens of derived attribute volumes. The amount of data available to the interpreter has grown faster than the time available to interpret it. As a result, only a fraction of acquired seismic data is typically examined in detail, and interpretation time is often spent building first-pass objects rather than testing geological hypotheses.

Machine learning has helped to address parts of this workflow. Deep-learning approaches have been applied to salt interpretation, fault detection, horizon tracking, relative geologic time estimation, and seismic facies classification. Wu et al. (2019) demonstrated that synthetic training data could be used to train a convolutional neural network for 3D seismic fault segmentation. Deep-learning approaches to relative geologic time have also shown that structural and stratigraphic interpretation can be supported by learned representations of seismic geometry (Geng et al., 2020; Bi et al., 2021). In seismic facies analysis, machine-learning methods have been used to classify seismic texture and reflection character more systematically (Wrona et al., 2018; Puzyrev and Elders, 2020).

These methods can reduce interpretation cycle time and improve repeatability, especially for well-defined tasks with sufficient labels. However, many applications remain narrow. They are trained for a specific objective, often on a specific dataset, and their generalisation to new regions depends strongly on the training distribution and label availability. This limitation has been highlighted in broader reviews of machine learning for seismic exploration, where generalisation, label scarcity, and workflow integration remain persistent challenges (Khosro Anjom et al., 2024).

Foundation models mark a shift in this paradigm. Rather than training a new model from scratch for each interpretation problem, a seismic foundation model (SFM) is first pretrained on large volumes of unlabelled seismic data. Through self-supervised learning, the model learns structural and stratigraphic patterns directly from seismic images. It can then be adapted to downstream tasks using smaller, task-specific datasets. The goal is not simply to automate more interpretation steps, but to

¹ TGS

* Corresponding author, E-mail: Henri.Houllevigue@tgs.com

encode reusable seismic representations that can support many interpretation workflows.

The foundation-model approach builds on advances in self-supervised learning, including masked autoencoders, in which models learn by reconstructing missing portions of an input from visible context (He et al., 2022). In geophysics, Sheng et al. (2023) introduced the concept of a seismic foundation model trained on unlabelled seismic images and adapted to multiple downstream tasks. Sansal et al. (2025) extended this work by demonstrating the scalability of a 3D seismic foundation model trained on a large global seismic dataset and fine-tuned for salt segmentation.

Salt was an appropriate first downstream task because it is geologically important, visually distinctive, and supported by high-quality interpretation labels. But salt is only one part of the interpretation problem. A useful geological interpretation workflow also requires faults, horizons, depositional elements, geobodies, facies, and well calibration. In this article, we focus on downstream applications that move the foundation-model approach closer to geological interpretation: structural discontinuities, depositional elements, and seismic facies.

From workstation to foundation-assisted interpretation

A conventional seismic interpretation workflow is organised around a series of connected tasks. The interpreter loads and quality-controls the data, ties wells to seismic, picks horizons, interprets faults, generates attributes, identifies geobodies, correlates well information, builds structural and stratigraphic models, and finally evaluates prospects and uncertainty. Each step depends on the previous one, but the workflow is rarely linear. Interpretation is iterative: a new fault interpretation may revise a horizon map, a channel interpretation may revise the depositional model, a well tie may require rethinking seismic phase, polarity, or time-depth relationships.

This workflow is also multiscale. Fault interpretation may require recognising local discontinuities on individual sections while also understanding regional structural style and variations in signal-to-noise ratio. Channel and geobody interpretation may require integrating subtle amplitude patterns, planform morphol-

ogy, stratigraphic position, and depositional context across a large area. Facies interpretation may require linking seismic texture to geological meaning using sparse well control. Experienced interpreters develop an internal model of these relationships over years of practice. They learn not only what a feature looks like but also how it behaves in different geological settings and seismic imaging conditions.

Interpretation is also tied to business scale. The level of detail required for frontier exploration differs from that required for reservoir characterisation. A useful interpretation workflow must therefore produce scale-coherent results: features should be resolved at a level appropriate to the decision being made. A regional exploration interpretation may focus on major structural elements, depositional fairways, and high-level prospect risk. A reservoir-scale interpretation may require identification of smaller faults, detailed geobody boundaries, and closer integration with wells.

Task-oriented machine learning has been integrated into this workflow as a point solution. A fault model may accelerate fault-stick generation. A salt model may provide an initial salt body. A facies model may classify seismic textures within a local interval. These tools can be powerful, but they usually do not share information. They do not naturally carry geological context from one task to another.

Foundation-assisted interpretation changes the model's role (Figure 1). The foundation model is not just another task-specific algorithm. It is a reusable representation engine. The same pretrained encoder can support multiple downstream tasks because it has already learnt a broad set of seismic patterns during self-supervised training. In this sense, the model begins to act less like an isolated automation tool and more like a seismic prior that can be adapted to different interpretation objectives.

However, the practical value of such a system depends on workflow integration. Rapid prediction alone is not enough. The interpreter must be able to edit, accept, reject, merge and evaluate model outputs efficiently. Without user-friendly editing and quality control, time saved during prediction may be lost later in the interpretation cycle. Foundation-assisted interpretation must therefore be designed as an interpreter-in-the-loop workflow, not as an isolated batch-inference exercise.

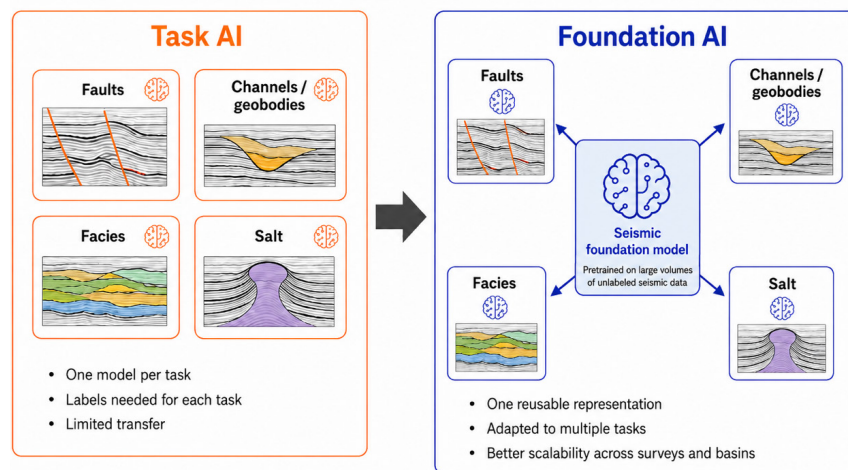


Figure 1 From task AI to foundation-assisted interpretation. Task AI uses separate models for individual interpretation problems, while foundation AI uses a reusable seismic representation that can be adapted across multiple tasks.

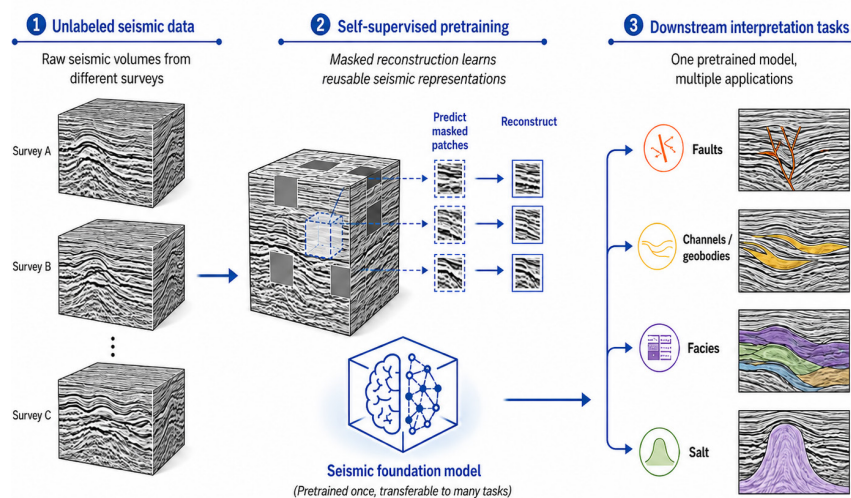


Figure 2 Self-supervised seismic foundation model concept. Unlabelled seismic volumes are used for masked reconstruction during pretraining, producing a reusable seismic representation that can be adapted to downstream interpretation tasks.

Seismic foundation model pretraining

The seismic foundation model used in this work is based on a 3D masked autoencoder architecture with a vision transformer backbone. During pretraining, seismic volumes are divided into 3D patches. A large percentage of these patches is masked, and the model is trained to reconstruct the missing seismic information from the visible context. This task forces the model to learn spatial continuity, reflector geometry, discontinuities, terminations, texture, and other seismic patterns without requiring manual labels.

This self-supervised formulation is well suited to seismic data. Seismic libraries contain large amounts of unlabelled data, whereas high-quality interpretation labels are expensive and unevenly distributed. By first learning from the seismic data itself, the model can leverage the full scale of the seismic archive before being fine-tuned for a specific interpretation task. The pretrained encoder can then be frozen or lightly adapted, while a smaller task-specific decoder is trained for segmentation, classification, interpolation, or other downstream applications (Figure 2).

The 3D formulation is important because interpretation is fundamentally volumetric. Faults, channels and depositional systems are not independent 2D images; they exhibit spatial continuity and geometry in three dimensions. A 3D model can use context from inline, crossline, and depth or time directions simultaneously. Larger context windows allow the model to capture broader geological patterns rather than only local texture.

Sansal et al. (2025) demonstrated this approach with a 3D ViT-MAE model trained on a global dataset of 63 seismic surveys. The model used masked reconstruction during pretraining and was then adapted for salt segmentation. The salt task showed strong generalisation compared with conventional task-specific training and established a framework for scaling seismic foundation models beyond single-dataset applications. That work also emphasised that large-scale model training depends on AI-ready seismic infrastructure, including cloud-native data access and the MDIO data format (Sansal et al., 2023; Sansal et al., 2025).

The question addressed here is how the same foundation-model paradigm can be extended from salt segmentation to a broader set of interpretation tasks. The examples are selected to illustrate distinct parts of the interpretation workflow rather than to present a single automated interpretation system.

Structural interpretation: faults, horizons and relative geologic time

Faults and horizons define the subsurface structural framework. They control trap geometry, compartmentalisation, migration pathways, seal risk, and uncertainty in structural maps. Errors in fault interpretation can propagate through horizon mapping, depth conversion, volumetrics, and prospect risk assessment.

Horizons also provide the building blocks for Relative Geologic Time (RGT). RGT is a continuous volume that assigns each seismic sample a relative chronostratigraphic position. Rather than mapping amplitude directly, RGT transforms the seismic image into a stratigraphic coordinate system, where each sample is tagged by its relative depositional order. RGT does not necessarily provide absolute geological age, but it can capture stratigraphic architecture within a continuous 3D framework.

Because RGT encodes relative depositional order, it is valuable for stratigraphic analysis, including onlap and offlap patterns, unconformities, channel systems, and depositional sequences. However, major faults disrupt the reflector continuity that horizon and RGT workflows rely on. For this reason, fault interpretation is often a prerequisite for building a consistent stratigraphic framework.

In the examples shown here, we focus on the fault component of this broader structural workflow. The current model output is a fault probability volume, not a horizon or RGT prediction. Nevertheless, the output is directly relevant to horizon and RGT interpretation because a consistent fault framework helps constrain where reflectors are offset, where horizons should be interrupted, and where stratigraphic continuity assumptions require careful handling.

Deep learning has already shown strong potential in structural interpretation. Wu et al. (2019) framed fault detection as a 3D image segmentation problem and trained an end-to-end convolutional neural network using synthetic seismic images and fault labels. Later work on relative geologic time showed that machine learning can also support structural interpretation by learning a volume that is consistent with seismic horizons and faults (Geng et al., 2020; Bi et al., 2021).

A foundation model offers a different form of support. Because the model has been pretrained to reconstruct seismic structure from sparse input patches, it learns representations that

are sensitive to reflector continuity, offsets, and terminations. When adapted to fault detection, the model can generate fault probability volumes that highlight discontinuities throughout the 3D volume (Figure 3). These volumes are not final fault surfaces. They are interpretation candidates that can be reviewed, filtered, connected, and converted into structural frameworks.

The value is not only speed. The value is consistency and coverage. A model can scan an entire survey, or multiple processing versions of a survey, using the same learnt criteria. This allows the interpreter to compare structural expressions across large

areas and focus time on geological QC: which faults are real, which are stratigraphic edges, which are processing artifacts, and which need to be incorporated into the structural model.

A practical workflow is to use the foundation model as a first-pass structural screening tool. The model generates a fault probability volume. The interpreter then selects a probability range and size or continuity threshold appropriate to the project objective. High-confidence, project-relevant faults can be accepted with limited editing, while ambiguous or lower-confidence features are reviewed manually. The interpreter then completes

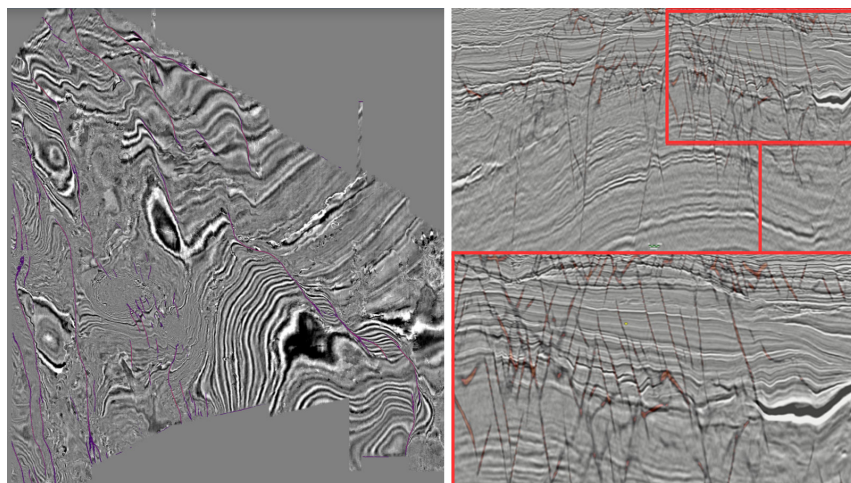


Figure 3 Fault probability responses on seismic sections from the open Parihaka 3D dataset (left) and a Brazil multi-client 3D dataset, with a zoomed view of the highlighted section (right). The foundation-model output highlights candidate structural discontinuities for interpreter review.

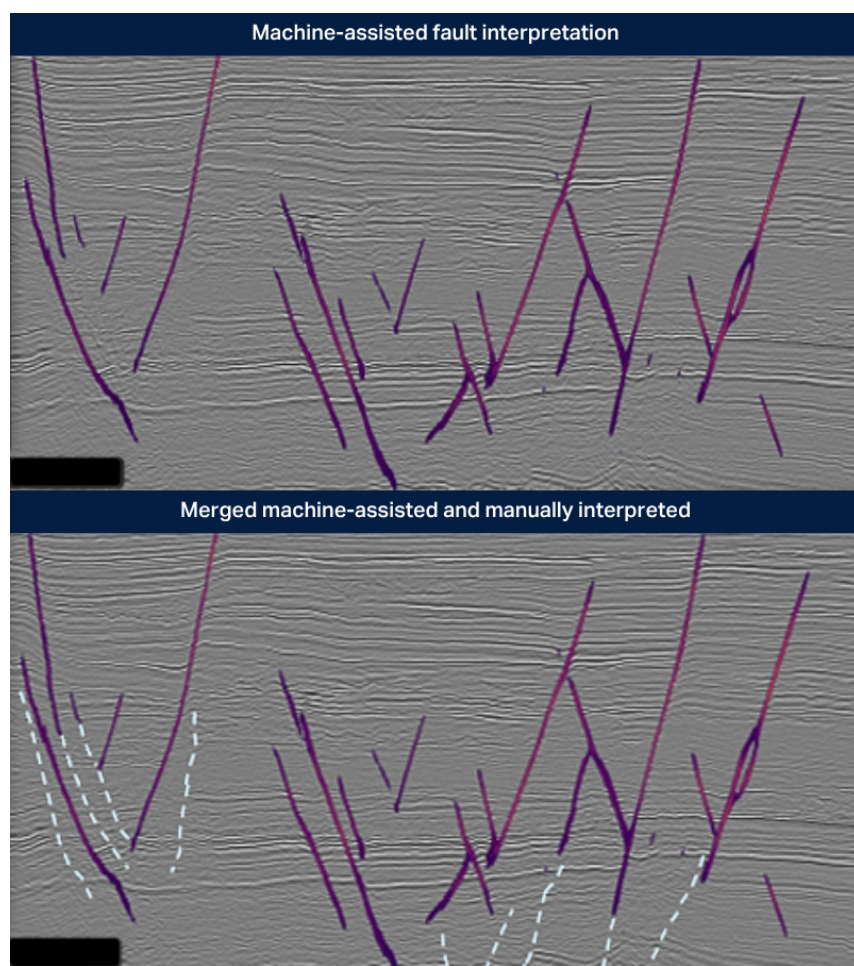


Figure 4 Interpreter-in-the-loop fault refinement. Foundation-model fault probabilities can identify high-probability discontinuities, while the interpreter adds, removes, or edits features that are important for the geological objective.

the residual fault interpretation required for the project and merges the machine-assisted and manually interpreted components into a final structural framework (Figure 4).

This distinction is important because not all faults have the same business relevance. In frontier exploration, the interpreter may be most concerned with major faults that define trap geometry, migration pathways, or regional structural style. In reservoir characterisation, smaller faults may become critical for compartmentalisation. A useful AI-assisted workflow must therefore allow the interpreter to filter and edit model outputs according to geological and business questions.

The same principle applies to horizons and RGT. A foundation-assisted structural workflow should not treat faults, horizons and stratigraphic ordering as independent products. Fault probability volumes can provide constraints for subsequent horizon tracking and RGT construction, while horizon and RGT consistency can, in turn, help to identify missing, overconnected, or geologically implausible faults. This joint review is where the broader value of a foundation-model representation becomes important, even when the first demonstrated output is fault probability.

Geological settings also matter. In many passive-margin settings, extensional and normal faults may produce clear reflector offsets, terminations and fault-plane geometries that are well suited to current 3D seismic interpretation models. Compressional and thrust systems can be more challenging because folds, ramps, flats and faulted folds may be ambiguous even for human interpreters. In these settings, model outputs should be treated as screening products that require careful structural QC and, where appropriate, integration with regional tectonic models.

Depositional elements and geobodies: mapping stratigraphic architecture

Channels, lobes, and other depositional geobodies pose a different set of interpretation challenges. Unlike faults, which are often expressed as discontinuities, depositional elements may appear

as subtle amplitude patterns, changes in seismic texture, tuning effects, or platform geometries visible only on selected slices or attributes. Interpreters commonly use spectral decomposition, sweetness, RMS amplitude, coherence, curvature and stratal slicing to enhance these features. The process is powerful but iterative and can require many attribute volumes and significant expert judgment.

A foundation model can learn depositional patterns directly from the seismic image. During pretraining, the model is exposed to many examples of reflector geometries, channelised forms, lobate geometries, truncations and stratigraphic organisations. When adapted to a channel or geobody task, the model can produce probability volumes that highlight candidate depositional elements across the 3D survey (Figure 5).

This does not eliminate the need for seismic stratigraphy. A high-probability channel response must still be interpreted in context: its stratigraphic position, relationship to surrounding reflectors, depositional direction, reservoir significance, and uncertainty. But it changes the starting point. Instead of manually searching through a large seismic volume and numerous attributes, the interpreter can begin with a model-generated candidate volume and use geological judgment to validate, edit, and rank the features.

This is particularly important for screening. Exploration teams often need to evaluate many surveys, many processing versions, or many intervals before deciding where to spend detailed interpretation time. A channel probability volume provides a repeatable first-pass view of depositional architecture. It helps to identify where potential reservoir elements may be present, where additional interpretation is needed, and where the seismic response is ambiguous.

The current results also underscore the importance of data quality and geological context. In areas with a high signal-to-noise ratio and clear depositional expression, the model can delineate channel boundaries and support rapid screening of depositional fairways. In more challenging environments, channel-like respons-

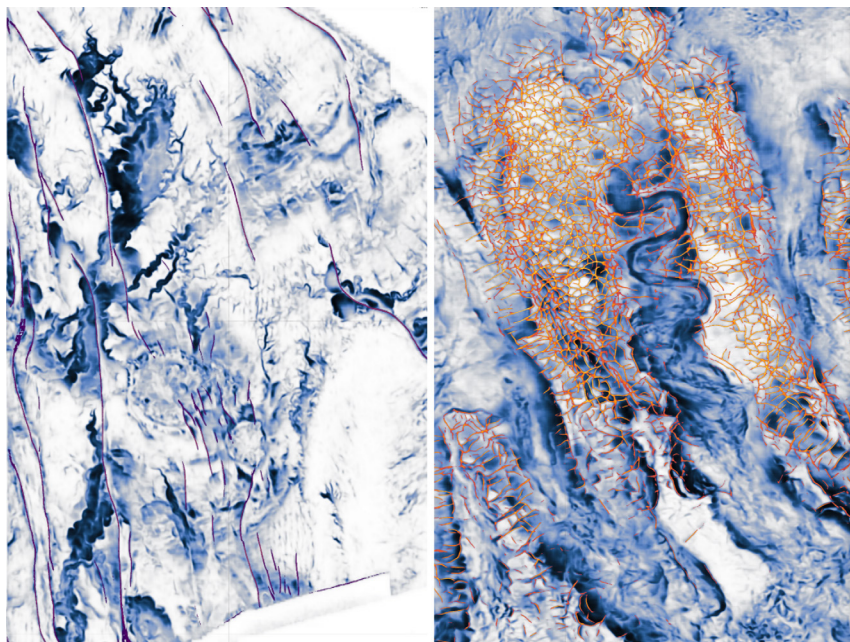


Figure 5 Fault and channel probability responses on horizontal slices from the open Parihaka 3D seismic dataset (left) and a Brazil multi-client 3D seismic dataset (right). The foundation-model predictions highlight candidate structural and depositional features for first-pass interpretation and screening.

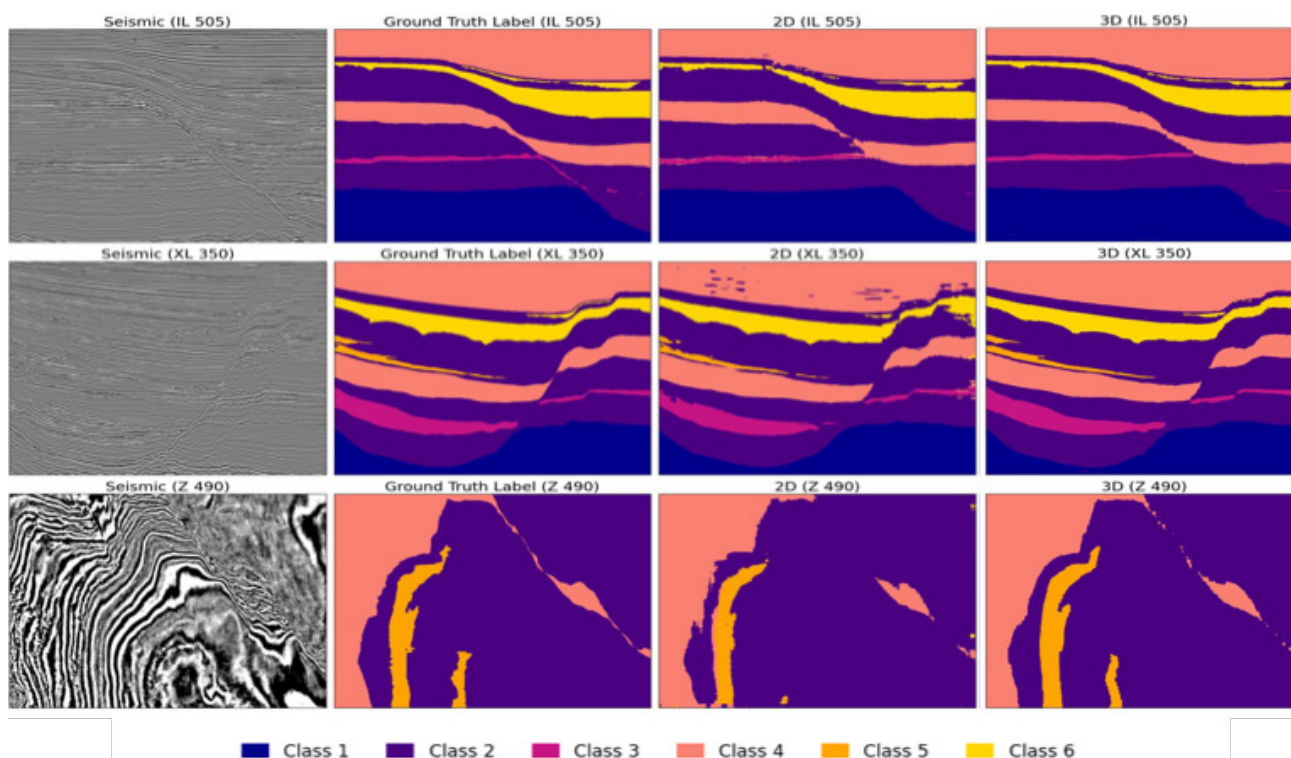


Figure 6 Seismic facies interpretation example from the open Parihaka 3D seismic dataset. The SFM-based model produces spatially coherent seismic facies predictions across inline, crossline, and depth views. The result should be interpreted as seismic facies, not direct lithology prediction.

es may require careful quality control to distinguish depositional features from faults, acquisition footprint, processing artifacts, or other stratigraphic discontinuities. This is not a failure of the interpreter-led workflow; it is a reminder that probability volumes must be reviewed within a geological context.

For this reason, channel and geobody interpretation is best treated as an assisted interpretation workflow. The model provides candidate geobodies and probability volumes. The interpreter validates them using seismic morphology, stratigraphic position, attributes, wells where available, and regional depositional understanding. The output is not simply a model classification; it is an interpreted depositional product supported by AI-generated candidates.

Seismic facies: from texture to geological meaning

Seismic facies interpretation connects seismic texture, amplitude, continuity and geometry to geological meaning. It is also one of the most difficult downstream tasks to define. Unlike salt or faults, facies labels are often interpretive, scale-dependent and tied to local geological contexts. A seismic facies class may represent a texture or reflection configuration rather than a unique lithology. Calibration to wells, cores, stratigraphic surfaces and regional geological understanding is therefore essential.

Machine learning has a strong history in seismic facies analysis. Wrona et al. (2018) demonstrated that supervised machine-learning methods can classify common seismic facies patterns and make facies analysis more systematic. Puzyrev and Elders (2020) explored unsupervised seismic facies classification using a deep convolutional autoencoder, showing how learned representations can support facies mapping without manually labelled examples.

The foundation-model approach builds on this idea but shifts the training strategy. A seismic foundation model does not directly observe rock type. It observes seismic patterns. When adapted for facies interpretation, the model learns to classify seismic expressions using the provided labels and geological framework. The resulting prediction can be useful, but it should be treated as a seismic facies or interpretation-facies product, not as an unconditional lithology model.

Figure 6 shows an example from the open Parihaka 3D seismic dataset in New Zealand, where expert-interpreted seismic facies provide a benchmark for evaluating facies prediction. In this example, the SFM-based facies model produces spatially coherent bodies across the seismic volume. This is important because facies interpretation is not only a local texture-classification problem. Interpreters also care about whether the predicted bodies are laterally consistent, stratigraphically plausible, and continuous across inline, crossline, and depth views.

The key point is not the numerical score alone, but the geological behaviour of the result. A useful facies prediction should help the interpreter to identify seismic texture domains, stratigraphic organisation and possible depositional patterns. It should also make areas of uncertainty easier to review. The model output is therefore best viewed as an interpretation aid: it can accelerate screening and improve consistency, but the geological meaning of each facies class still needs to be calibrated against available wells, regional geology and interpreter judgment.

Discussion: interpreter-led foundation AI

The main impact of seismic foundation models is not full automation. It is a shift in how interpreter time is spent. In conventional workflows, geoscientists spend significant effort

generating first-pass objects: scanning volumes, picking faults, testing attributes, identifying geobodies and classifying seismic textures. In a foundation-assisted workflow, the model can repeatedly and consistently generate candidate probability volumes, allowing interpreters to focus more on validation, integration and geological decision-making.

This distinction is important because model outputs remain probabilistic. They depend on data quality, training labels and geological context, and should be treated as interpretation candidates rather than ground truth. They can reveal patterns, accelerate screening and improve repeatability, but they do not carry geological meaning on their own. Interpreters remain responsible for quality control, calibration to wells, uncertainty assessment, and final geological interpretation.

Generalisation should also be validated rather than assumed. Foundation models are designed to generalise better than narrow, task-specific models, but local calibration or fine-tuning may still be needed. Facies prediction requires particular care because seismic facies are not necessarily lithofacies; their geological meaning depends on the labelling strategy, well control, stratigraphic framework, and depositional model.

The broader opportunity is to move beyond isolated task AI. Structural discontinuities, depositional elements, and facies can be generated as separate downstream products, but their value increases when they are reviewed alongside horizons, wells, attributes, and regional geological knowledge. This requires not only accurate models but also practical tools for visualisation, editing, uncertainty review, and export to interpretation platforms.

Conclusions

Seismic foundation models move interpretation beyond isolated task automation. By learning reusable representations from large volumes of unlabelled seismic data, they provide a common starting point for structural interpretation, detection of depositional elements, and seismic facies classification. The examples discussed here show that foundation-model outputs are best treated as probability volumes and interpretation candidates, not final geological answers.

Their value is to accelerate first-pass screening, improve repeatability and help interpreters focus on geological meaning. The interpreter remains central. Model outputs still require QC, calibration against wells and regional geology, and uncertainty assessment. The practical promise is therefore not automated interpretation but foundation-assisted interpretation: scaling the repetitive parts of interpretation while preserving human control over geological integration.

References

- Bi, Z., Wu, X., Geng, Z. and Li, H. [2021] Deep Relative Geologic Time: A Deep Learning Method for Simultaneously Interpreting 3-D Seismic Horizons and Faults. *Journal of Geophysical Research: Solid Earth*, **126**(9), e2021JB021882.
- Geng, Z., Wu, X., Shi, Y. and Fomel, S. [2020] Deep learning for relative geologic time and seismic horizons. *Geophysics*, **85**(4), WA87-WA100.
- He, K., Chen, X., Xie, S., Li, Y., Dollár, P. and Girshick, R. [2022] Masked Autoencoders Are Scalable Vision Learners. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*.
- Khosro Anjom, F. et al. [2024] Machine-learning for seismic exploration: Where are we and how far are we from the holy grail? *Geophysics*, **89**(1), WA157-WA178.
- Puzirev, V. and Elders, C. [2020] Unsupervised seismic facies classification using deep convolutional autoencoder. arXiv:2008.01995.
- Sansal, A., Kainkaryam, S., Lassecock, B. and Valenciano, A. [2023] MDIO: Open-source format for multidimensional energy data. *The Leading Edge*, **42**(7), 465-470.
- Sansal, A., Lassecock, B. and Valenciano, A. [2025] Scaling seismic foundation models. *First Break*, **43**(2), 69-74.
- Sheng, H., Wu, X., Si, X., Li, J., Zhang, S. and Duan, X. [2023] Seismic Foundation Model: A new generation deep learning model in geophysics. arXiv:2309.02791.
- Wrona, T., Pan, I., Gawthorpe, R. L. and Fossen, H. [2018] Seismic facies analysis using machine learning. *Geophysics*, **83**(5), O83-O95.
- Wu, X., Liang, L., Shi, Y. and Fomel, S. [2019] FaultSeg3D: Using synthetic data sets to train an end-to-end convolutional neural network for 3D seismic fault segmentation. *Geophysics*, **84**(3), IM35-IM45.