

UNIVERSITY OF OKLAHOMA

GRADUATE COLLEGE

MAKING SENSE OF MULTI-ATTRIBUTE MACHINE LEARNING:
GAINING SEISMIC INTERPRETATION INSIGHTS WITH SHAP FOR UNSUPERVISED
CLUSTERING

A THESIS

SUBMITTED TO THE GRADUATE FACULTY

in partial fulfillment of the requirements for the

Degree of

Master of Science

By HILMI PUTRA

Norman, Oklahoma

2026

MAKING SENSE OF MULTI-ATTRIBUTE MACHINE LEARNING:
GAINING SEISMIC INTERPRETATION INSIGHTS WITH SHAP FOR UNSUPERVISED
CLUSTERING

A THESIS APPROVED FOR THE
SCHOOL OF GEOSCIENCES

BY THE COMMITTEE CONSISTING OF

Dr. Heather Bedle, Chair

Dr. Hao Hu

Dr. David Lubo-Robles

ABSTRACT

In seismic multi-attribute analysis, unsupervised machine learning helps identify patterns inherent in the data across the multiple attribute dimensions. It summarizes the information intuitively for appraisal of the underlying geological objects responsible for the different seismic signal expressions. Conventionally, though, unsupervised machine learning products—in the form of volumes and maps—serves mainly as leads for the seismic interpreter, while applying further seismic attribute theory for interpretation still requires going back to the original input attribute volumes. This points to a deeper challenge: unsupervised methods inherently lack an interpretability layer, leaving the model's internal reasoning opaque to the interpreter.

Explainability frameworks such as SHapley Additive exPlanations (SHAP) have gained traction in supervised learning, where a well-defined output metric—typically a prediction probability or classification score—enables the attribution of input feature contributions. Without an equivalent metric, unsupervised models have largely been outside the scope of SHAP-based interpretability. To bridge this gap, this study utilizes probability distributions from a Gaussian Mixture Model (GMM) as a proxy metric for cluster membership confidence, effectively applying SHAP to an unsupervised multi-attribute machine learning workflow.

This integration enables analysis into how individual seismic attributes contribute to each cluster assignment. Which means each appraised cluster is then accompanied by geophysically meaningful information that highlights the key signal expressions and attributes characterizing it. This added interpretability works to equip the geoscientist with data-driven quantitative means to support their interpretation, while simultaneously grounding the machine learning results back to seismic attribute theory.

In this study, we found that the SHAP graphical descriptions align with seismic attribute characterization of geological targets in the Taranaki Basin, showing how the cluster is expressed in each attribute space. The graphical descriptions enhance the role of the unsupervised machine learning workflows in seismic multi-attribute analysis and provide intuition-driven mechanisms for quality control and trust in machine learning products.

TABLE OF CONTENTS

ABSTRACT.....	iv
TABLE OF CONTENTS	v
TABLE OF FIGURES	vi
1 Introduction.....	1
1.1 Data	3
1.2 Geological Target	6
2 Methodology	8
2.1 Seismic Attributes	10
2.1.1 Amplitude Attributes.....	11
2.1.2 Spectral Attributes.....	12
2.1.3 Structural Attributes	13
2.1.4 Texture Attributes.....	15
2.2 Gaussian Mixture Models (GMM)	16
2.2.1 Bayesian Information Criterion	18
2.3 SHapley Additive exPlanations (SHAP).....	20
3 Results and Discussion	24
3.1 GMM.....	24
3.2 SHAP Analysis.....	28
3.2.1 Cluster Characteristics	31
3.2.2 Attribute Contribution Performance	37
4 Conclusion	39
4.1 Limitations	40
4.2 Recommendations.....	40
5 References.....	42
Appendix.....	44

TABLE OF FIGURES

Figure 1.1 Map showing the location of Kokako 3D in New Zealand	3
<i>Figure 1.2 average spectrum (left), and a section showing the seabed reflection (right) of Kokako 3D Survey.</i>	4
Figure 1.3 Inline 1400. Stacking edge artifacts can be seen on the left, and a bit on the right edge. Reference Horizon is shown as black line	4
Figure 1.4 Time slice 1642ms. The location of the two available wells is marked as black dots. Blue polygon is the crop boundary that is used for this study area.	5
<i>Figure 1.5 Self Organizing Map (SOM) of the cropped preliminary volume for a visual differentiation of the seismic facies within (Kohonen, 1982). We can see the channels and fan structure on the time slice (left) and crossline sections (right).</i>	5
<i>Figure 1.6 A geological model for the morphology of a clinoform-slope and its common features through the Highstand and Lowstand Systems Tract. Similar to what we observe in the Neogene depositional of the Kokako 3D survey. from (Weimer, et al., 2007)</i>	6
Figure 1.7 Arbitrary line showing the clinoform successions. Colored lines are interpreted horizons corresponding to each phase of the stratigraphic sequence that highlights the different modes of progradation as consequence of the interplay between accommodation space and sediment supply. Bright red (lowest lines) represents the initial FSST (Falling Stage Systems Tract) of the interval of interest. Bright blue is LST phase. Green is the retrogradational TST (Transgressive Systems Tract). Dark blue is HST phase. Followed by dark red which is the final FSST that makes up the prominent feature, Giant Foresets Formation.	7
Figure 1.8 Clinoform breakpoint trajectory and interpretation of the basin's depositional evolution modified from Franzel and Back (2018). The colored numbers correspond to the phases' age in million years ago. The left diagram illustrates the successions that we can observe in Kokako 3D.	8
Figure 2.1 Study workflow	9
Table 2.2 attributes selected for this study and the signatures they are extracting	10
Figure 2.3 Principal Component Analysis eigenvectors variability across various amplitude-type attributes.....	11
Figure 2.4 envelope (horizon slice 284ms from reference horizon).....	12
<i>Figure 2.5 Average spectrum from the zone of interest. Blue is the original Kokako3D data, orange is after spectrum balancing is applied. Red lines are boundaries of the frequency bins for the spectral decomposition</i>	13
<i>Figure 2.6 RGB Blend map of 20.5, 35.5, 58Hz frequency magnitudes along with the PCA eigenvector values across the frequency bins to show the variability that the three selected bins represent.</i>	13
Figure 2.7 k1 x k2 (corendering) (horizon slice 284ms from reference horizon).....	14
Figure 2.8 outer product similarity (horizon slice 284ms from reference horizon).....	15
Figure 2.9 GLCM entropy (horizon slice 268ms from reference horizon).....	16

<i>Figure 2.10 An example of a Gaussian normal distribution (probability density function) with a mean value of 0 and standard deviation of 1(left). An example 2D Gaussian distribution of data points with a skewed oval direction indicating a significant off-diagonal element of the covariance matrix(right)</i>	17
Figure 2.11 How to pick optimum cluster number for GMM using BIC. Top shows a smaller dataset where the BIC value has a lowest point which is at the optimum number of cluster. Bottom is on a larger dataset where the BIC value keeps improving (decreasing) so the optimum cluster number is interpreted at the elbow where the improvement converges. The right graph is the rate of decrease to make the elbow appear more obvious.....	19
Figure 2.12 How SHAP works by creating coalitions and producing the outputs with and without the feature being analyzed. Then, it weights the result by the number of coalition size (in this example 4) and how many possible coalitions for each size so that each coalition size has the same weight towards the final SHAP value.	20
Figure 2.13 How SHAP works and how its logic can be applied to unsupervised ML with GMM.	21
Figure 2.14 SHAP waterfall plot. This graph illustrates what happens behind every datapoint: each attribute input will have a contribution—SHAP value—that adds up towards the final probability value/degree of membership to the cluster. Here we see one of the datapoint that is a member of cluster 4 that is mainly pushed into membership (to the right) by the low relative value k2 (most negative curvature), and low value glcm entropy second.	22
Figure 3.1 envelope (left) and GLCM entropy (right) attribute histogram distribution with the cluster Gaussian distributions overlaid. The combination of cluster Gaussian normal distributions (colored lines) recreate the original non-normal distribution (grey graph), taking into account the juxtaposition of cluster characteristics across attribute values. All histogram graphs can be seen on appendix A.1	24
Figure 3.2 Submarine fan architecture (red/top) seen with the straight slope channels feeding into the unconfined fan deposit. Sinuous channelized flow (blue/right) seen with the bifurcation of channel—yellow facies—with overbank facies surrounding it.	25
Figure 3.3 The first progradational-upbuilding HST phase where the large straight channels initiates at the clinoform breakpoint and concludes into fan structures at the slope landing. White dashed line show the interpreted clinoform breakpoint and the red dashed line for the interpreted slope landing into the basin floor.....	26
Figure 3.4 The end of the first progradational-upbuilding LST phase (top), later overlain by the seismic gap of the retrogradational TST phase. Then the second progradational-upbuilding HST phase (bottom). Red dashed lines show the slope transition where white dashed lines show the interpreted clinoform breakpoint. The progradation characteristic of this phase can be seen with the shelf landing is not as distinct moving North Westward (red arrows) indicating a gentler slope transition.	27
Figure 3.5 the final progradational-downbuilding phase (Giant Foresets Formation) where it's initiated with the mass transport complex and features the distributary feeder channels flowing	

into the slope gulleys indicating its closeness to shoreline and erosional surface (interpreted with solid white line) from the forced regression	28
Figure 3.6 Examples of cluster attribute behavior illustrated by global SHAP summary plot: beeswarm plot. The dots represent each datapoint from the observed dataset, its x-axis is the SHAP value of the corresponding attribute, vertically arranged by attribute importance, and the color represents the relative attribute value.	29
Figure 3.7 k2 (most negative curvature) contribution towards cluster 7 plotted against its normalized value shows the complex behavior of the attribute characteristic in identifying this cluster. We find a generally supporting low value (blue, pushing towards membership—to the right), but the direction changes when moving into the extreme low attribute value. The resulting SHAP beeswarm plot is shown at the top.	30
Figure 3.8 Cluster 4: Channel (low entropy) body, mainly characterized by the geometrical input attributes k1 and k2. The high magnitude of low frequencies (20 and 35Hz) supporting the membership while high 58Hz magnitude pulls away from membership further indicate the tuning of vertical layering of the channel structure to the lower frequency bins. Another interesting note is GLCM entropy can be seen supporting with its low values, which means this facies have a well layered, undisrupted, reflection.....	31
Figure 3.9 Cluster 7: Channel overbank or edges (thin layers). While having also K1 and K2 as the main indication similar to Cluster 4, high magnitude of the low frequencies are negatively supporting instead (pulling away from membership) while intermediate high magnitude of high frequency (58Hz) is pushing towards membership, indicating a tighter reflection layering.....	33
Figure 3.10 Cluster 8: Fan facies, mainly indicated by high reflection energy across the spectrum. GLCM entropy hinting a positive (pushing) contribution from the low values meaning a generally coherent, unbroken reflection spatially. As expected of a submarine fan deposit.	34
Figure 3.11 Xline 2693 of the GMM cluster output corendered with the seismic trace (peaks), showing the location of well Toutouwai-1. We can see how the well penetrates through some of the key architectural feature clusters. 108ms from reference horizon (top red dashed line), on the well location, the top white arrow points at where the GR log (white wiggle) is low (dips left) which coincides with the channel facies (cluster 4) even when P-wave (colored well column) doesn't seem to exhibit any clear distinction. The bottom white arrow points at another anomaly this time indicated by a low velocity (green-yellow) and high GR which happens to coincide with the overbank facies (cluster 7) on the side of a channel facies (cluster 4).	35
Figure 3.12 Xline 3057 of the GMM cluster output corendered with the seismic trace (peaks), showing the location of well Takapou-1. The large fan lobe (cluster 8) at 284ms from reference horizon (the lower red dashed line) can be seen just to the left (NE) of the well location, while the well itself seems to go through the complex mix of facies with GR highs and lows (black wiggle inside the well column) and velocity anomalies (marked with the two lower arrows) indicating a complex layering of different lithologies. The top arrow points at the mass transport complex channel-like structure at about -148ms coincide with a slightly high p-wave velocity on the well log and marked as channel facies (cluster 4).	36

Figure 3.13 Cluster 5 of the GMM ran with 13Hz spectral magnitude as one of the inputs.
Horizon slices show how cluster 5 does not exhibit a distinct meaningful signature even though
showing a high contribution from 13Hz magnitude. 37

1 Introduction

In exploration seismology, specifically reflection seismology, we use the echoes of sound waves to observe and characterize the subsurface. A seismic source is excited, typically in the form of an impulse (i.e. a detonation) or a specifically engineered signal, to produce a wave that propagates into the subsurface. Portions of this energy are reflected whenever the wave encounters impedance contrasts, change of density and velocity, typically representing interfaces between different rock units. These reflected waves are then picked up by receivers (geophones or hydrophones) and recorded into the seismic traces (seismogram). The manner in which the seismic wave is reflected and attenuated will be evident in these recordings—analogue to how sound reflections differ between a padded studio room and a concert hall. In this sense, analyzing the recorded reflection signals allows us to infer the characteristics of the objects or layers that reflect them.

We apply various mathematical operators to the seismic traces to measure and quantify the features of the recorded signals. These extracted values, which we call seismic attributes, capture features such as amplitude, spectrum content, and spatial arrangement. They represent the different mechanisms of how subsurface properties influence specific characteristics of the seismic response that may relate to variations in lithology, fluid content, or depositional patterns (Chopra, et al., 2024). By quantifying these signal expressions, seismic attributes enhance the interpretability of seismic data and aid geoscientists in identifying and characterizing subsurface features of interest.

Analyzing multiple seismic attributes together provide complementary perspectives towards the seismic data, each emphasizing different expressions of the subsurface. This multi-attribute analysis allows geophysicists to delineate subtle anomalies in the seismic response that becomes apparent when viewed side-by-side (Zhao, et al., 2016; La Marca, et al., 2022). However, the increasing number of available attributes introduces complexity, as each attribute represents one dimension in a growing multi-dimensional dataset while an interpreter would effectively only be engaged with one of them at a time. Machine learning provides a means to capture this complexity by recognizing patterns across multiple attribute dimensions and summarizing them in a more meaningful, efficient, and objective way.

Unsupervised clustering methods group data points based on similarity or continuity in values—quantified in terms of a probability or distance metric in the multidimensional space—with the goal

of capturing the underlying geological facies that give rise to those **shared** signal expressions (Coléou, et al., 2003). Although optimized clustering workflows can effectively reveal and visualize patterns across the multi-attribute seismic data (La Marca, et al., 2024), interpretation of unsupervised machine learning results is often carried out separately from the conventional multi-attribute seismic facies analysis. In practice, the machine learning output is typically treated as a lead—an indicator of where to look—rather than its own supporting evidence for geophysical interpretation. The resulting clusters are visually compared to seismic sections or attribute maps, and their meanings are inferred by associating them with the characteristics observed in the highlighted zones. This creates a disconnect between the data-driven and intuition-driven components of interpretation.

To address this, we introduce the use of SHapley Additive exPlanations (SHAP). SHAP provides a consistent and theoretically grounded framework for quantifying the contribution of individual input features to a machine learning model’s output, (Lundberg, et al., 2017) in this case, clustering probabilities. Since SHAP operates by observing changes in output probability values, it is conventionally applied to supervised models (probability of label assignment). However, as an unsupervised clustering model, Gaussian Mixture Models (GMMs) are uniquely suited for this integration, as they produce probability distributions over cluster assignments, which can be directly used as the output that SHAP evaluates. By applying SHAP to unsupervised GMM clustering, we uncover what drives each clustering assignment—how it identifies the group—making the patterns not only observable but also *explainable*.

The machine learning results are inherently a representation of what is fed into the model: subsurface expressions captured by the input seismic attributes. The added context, thus, provide the link between these expressions to the seismic attributes that define them, complementing the clustering results with the geophysical information that an interpreter can apply intuition to when characterizing what each cluster represents. Furthermore, the analysis would reveal how *well* the selected attributes perform in differentiating the patterns or any potential underperformance (Lubo-Robles, et al., 2020). So, it can be used as means to improve attribute selection for machine learning model input or as an added caution for the interpreter.

With this intent, we test the GMM-SHAP approach to the Kokako 3D seismic dataset targeting the Neogene clinoform architectural elements to interpret the basin’s sedimentary succession and

acquire its corresponding seismic attribute characteristics from the machine learning with more confidence.

1.1 Data

The Kokako 3D marine seismic survey is located in the Kokako Field within the Tasman Sea, offshore Taranaki, New Zealand, just west of the Taranaki Peninsula (Figure 1.1). The survey was acquired by WesternGeco (WGC) using the M/V Western Monarch from 2–21 April 2013. Data acquisition employed a system of two seismic sources and ten streamers across 39 sail lines, covering an area of approximately 593 square kilometers. The dataset has an inline trace spacing of 25 m and a crossline spacing of 12.5 m.

This project utilizes the time domain PSDM full-angle-stack volume. We note that this dataset has a positive standard (normal) polarity zero phase wavelet, where a peak corresponds to an increase in acoustic impedance, as verified by seabed reflection (Figure 1.2). The processing steps were reported to include Structurally Oriented Filtering (SOF) and MigQ by DownUnder GeoSolutions which significantly improve relevant signal and image clarity crucial for our seismic attributes computation (2013). From analysis of the spectrum, we note that the dataset has a peak frequency of 44 Hz, corresponding to a dominant wavelet resolution of about 11.5 ms, with a usable bandwidth ranging from approximately 8 Hz to 78 Hz (Figure 1.2).

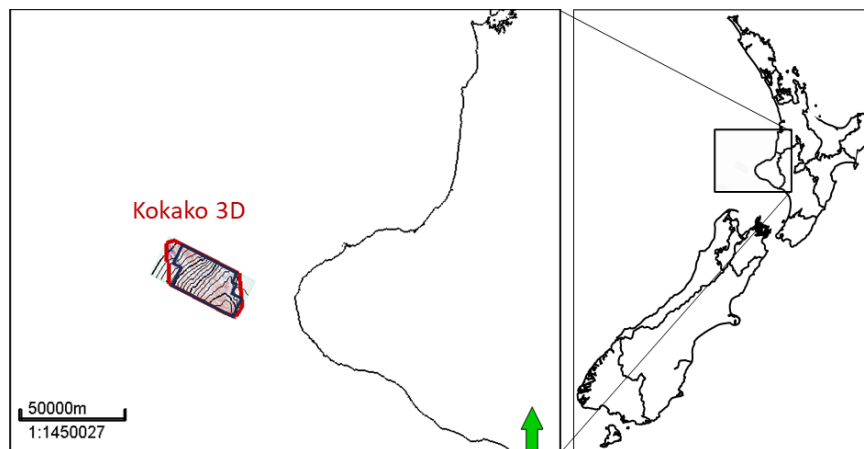


Figure 1.1 Map showing the location of Kokako 3D in New Zealand

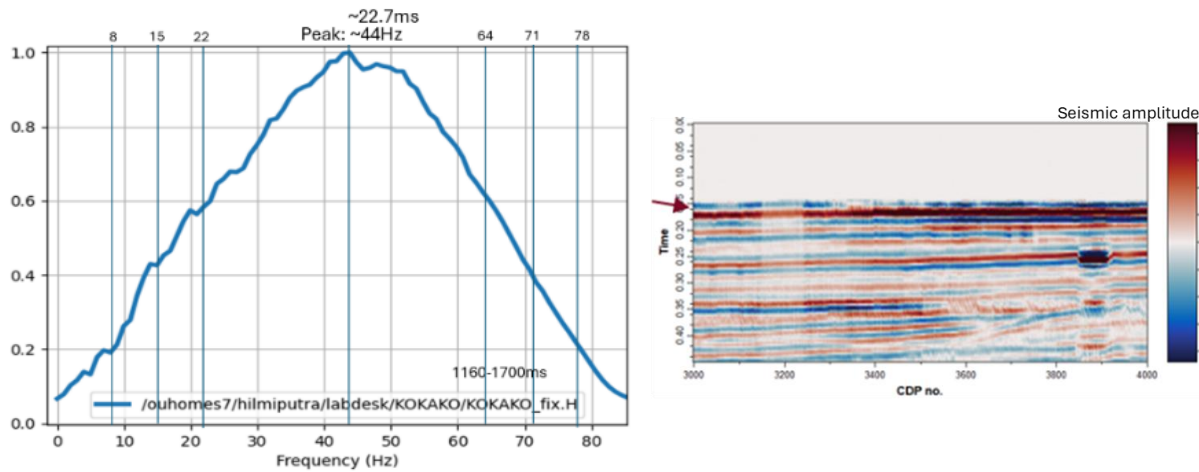


Figure 1.2 average spectrum (left), and a section showing the seabed reflection (right) of Kokako 3D Survey.

For this project, the edges of the seismic cube were cropped to remove stacking edge artifacts to ensure stable clustering performance, as these artifacts can introduce irrelevant attribute values during calculation (Figure 1.3). The geological objects of interest—Neogene clinoforms and Mass Transport Deposits—are within approximately 1100–2700ms, with the main target intervals situated at 1164–1772ms. Due to the tilted geometry inherent to clinoform deposition and subsequent tectonic deformation, observations are conducted along pseudo-horizon slices extracted relative to a reference horizon, quickly interpreted for one of the intermediate successions (2nd Upbuilding Phase, Unit 7; discussed in the Geological Target section), rather than along flat time slices, to preserve stratigraphic context and improve clarity of visualization (Figure 1.3).

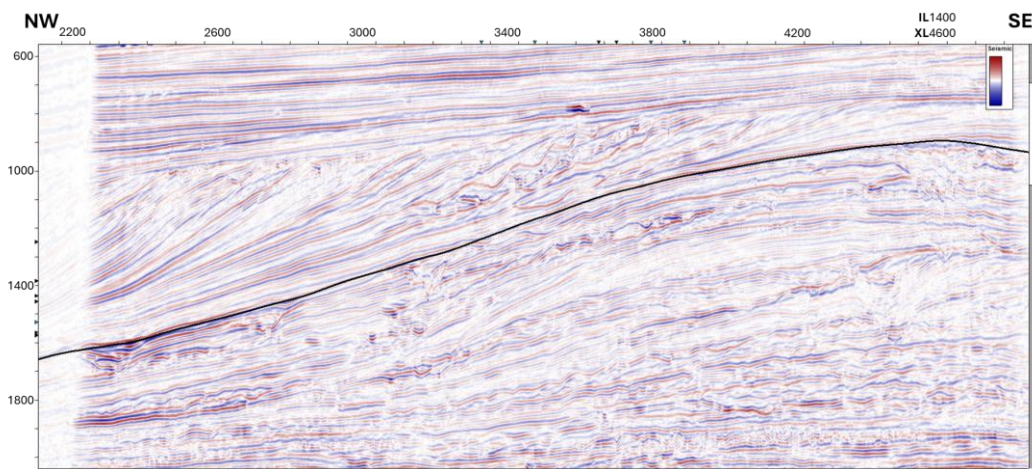


Figure 1.3 Inline 1400. Stacking edge artifacts can be seen on the left, and a bit on the right edge. Reference Horizon is shown as black line

Two wells are situated in the seismic survey and are in the vicinity of the targeted geology. Toutouwai-1 and Takapou-1 locations are shown on Figure 1.4. Both are located at the western corner of the cube. These wells are used as reference for lithological characteristics of our appraised seismic facies clusters or anomalies.

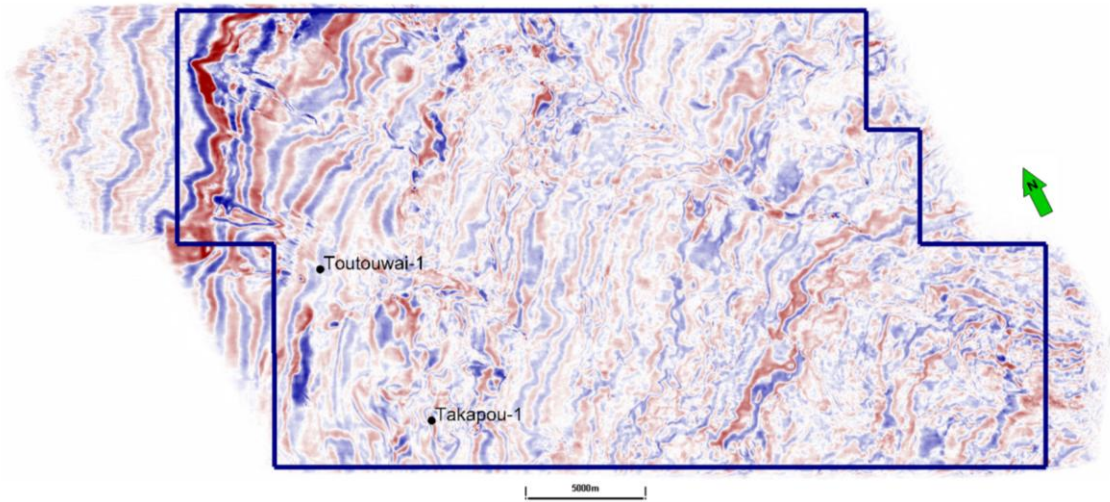


Figure 1.4 Time slice 1642ms. The location of the two available wells is marked as black dots. Blue polygon is the crop boundary that is used for this study area.

A reduced crop-volume was used for algorithm development and testing of the implementation workflow (see appendix A.5), spanning 1248–1600 ms two-way travel time (TWT), and bounded by crosslines 2940–3819 and inlines 1424 – 1730 (Figure 1.5). This crop-volume captures one of the interpreted submarine fan lobe structures within the study interval.

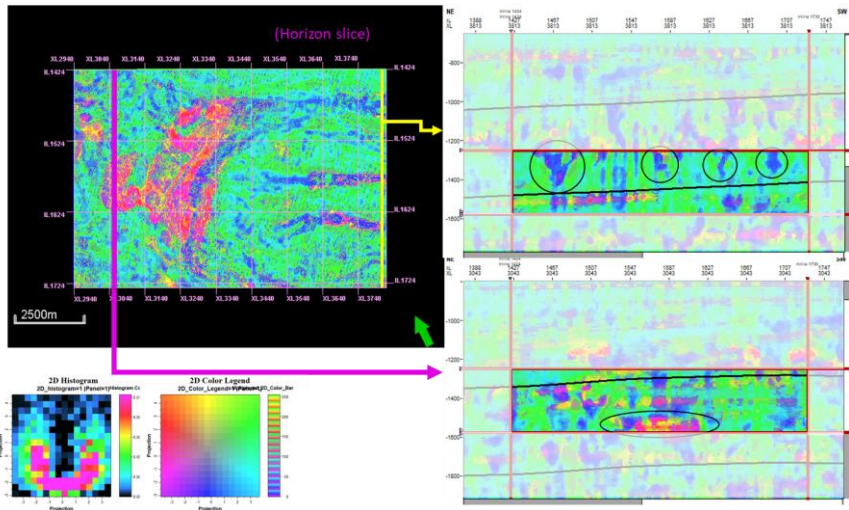


Figure 1.5 Self Organizing Map (SOM) of the cropped preliminary volume for a visual differentiation of the seismic facies within (Kohonen, 1982). We can see the channels and fan structure on the time slice (left) and crossline sections (right).

1.2 Geological Target

For validation of the machine learning-derived seismic facies interpretations, we refer to established deepwater mass transport geology and the expected seismic attribute responses of its characteristic architectural elements.

The Kokako 3D survey area mainly overlooks the distal slope region of the central Taranaki Basin, imaging a complex of deepwater mass transport deposits (MTDs) and turbidites with related channel-levee and submarine fan systems—architectural features commonly associated with this depositional environment.

Within this setting, key geomorphological features of interest include straight slope channels that initiate near the clinoform breakpoint, sinuous channels on less steeper sections along with its associated levees, submarine fan deposits, and mass transport complexes (Figure 1.6). These represent the middle to lower fan—lower slope to basin floor—components of the Neogene clinoform system which record repeated phases of progradation – retrogradation with mass wasting along the clinoform breakpoint in the central Taranaki Basin, as documented by Franzel and Back (2018; 2019).

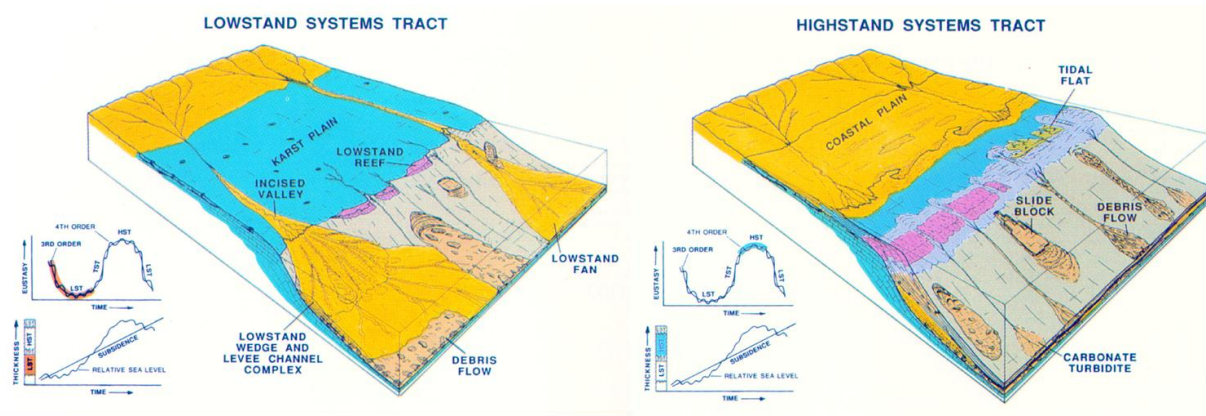


Figure 1.6 A geological model for the morphology of a clinoform-slope and its common features through the Highstand and Lowstand Systems Tract. Similar to what we observe in the Neogene depositional of the Kokako 3D survey. from (Weimer, et al., 2007)

These architectural features reflect the clinoform trajectory development which is ultimately governed by the basin's tectono-stratigraphic evolution. The spectacular Miocene to Plio-Pleistocene clinoform succession of the central Taranaki Basin reflects the shifts in relative sea level and

sediment supply, representing the interplay between accommodation creation and sedimentation processes.

The work by Franzel and Back (2018) utilized RGB blends of frequency magnitudes (10, 20, and 30 Hz) occasionally combined with RMS amplitude maps for manual interpretation—the conventional human approach to facies interpretation. They concluded that evidence of significant mass wasting events could be identified within the study area that corresponds to each sequence-stratigraphic phases (Figure 1.8), in the order: downbuilding-upbuilding-retrogradation-upbuilding-downbuilding (corresponding to FSST-LST-TST-HST-FSST sequence stratigraphy phases). Upbuilding progradational units (lowstand and highstand system tracts) tend to develop larger deep-water channels, of which two such events were captured within the Kokako 3D study interval.

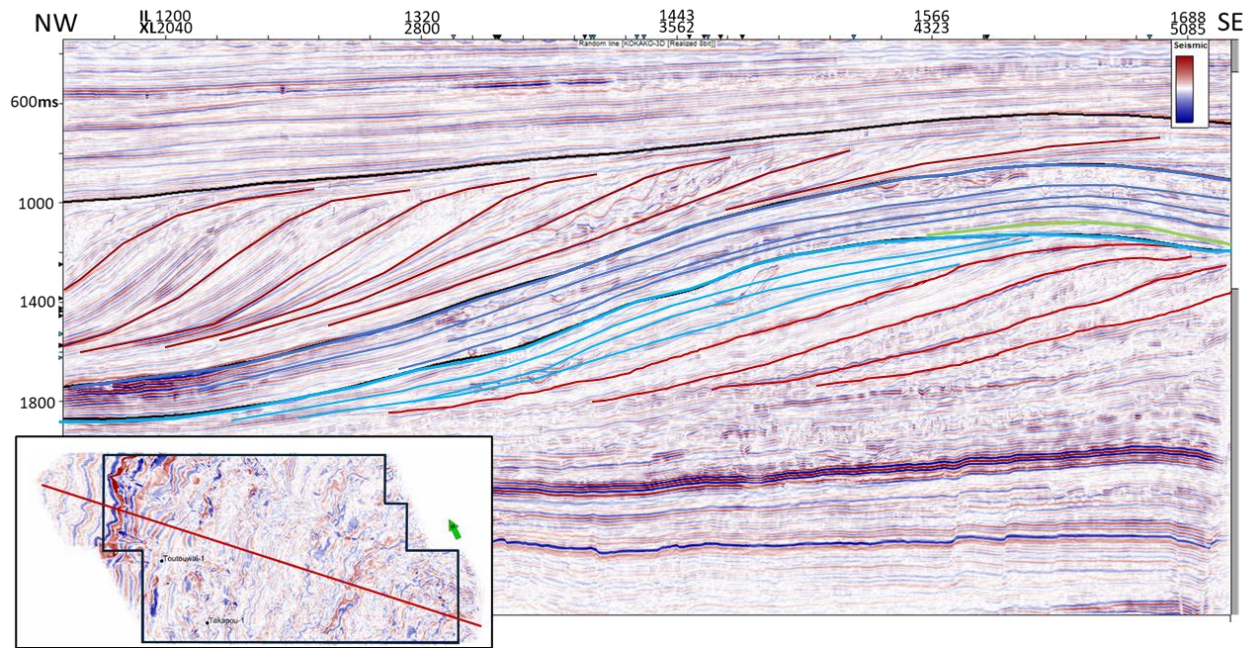
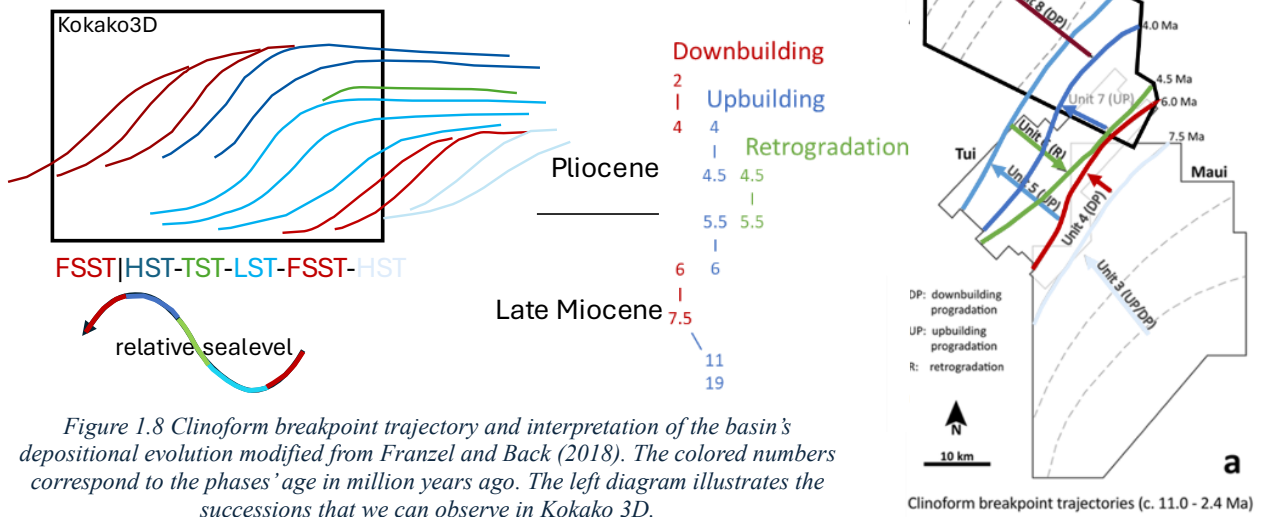


Figure 1.7 Arbitrary line showing the clinoform successions. Colored lines are interpreted horizons corresponding to each phase of the stratigraphic sequence that highlights the different modes of progradation as consequence of the interplay between accommodation space and sediment supply. Bright red (lowest lines) represents the initial FSST (Falling Stage Systems Tract) of the interval of interest. Bright blue is LST phase. Green is the retrogradational TST (Transgressive Systems Tract). Dark blue is HST phase. Followed by dark red which is the final FSST that makes up the prominent feature, Giant Foresets Formation.



The combination of the diverse and complex geological architecture—an array of deep water depositional elements and mass transport features—with the coverage by high-quality, industry-grade 3D seismic data, makes the Neogene clinoform succession of the Kokako 3D dataset an exceptional case study for evaluating multi-attribute machine learning approaches to seismic facies characterization.

2 Methodology

The objective of this study is to produce interpretable multi-attribute seismic facies by data-driven means, without supervision or prior interpretation. To do this, we perform a quantitative based evaluation of seismic expressions using mathematical operators, bringing forth the relevant geologic information, to which an explainable AI technique is applied to gain insight from the process for facies characterization. Thus, the approach consists of three main components (introduced in Figure 2.1): (1) seismic attributes generation and selection based on the relevant signal expressions, (2) ML clustering using the Gaussian Mixture Model (GMM), and (3) interpretation enhanced by SHAP analysis.

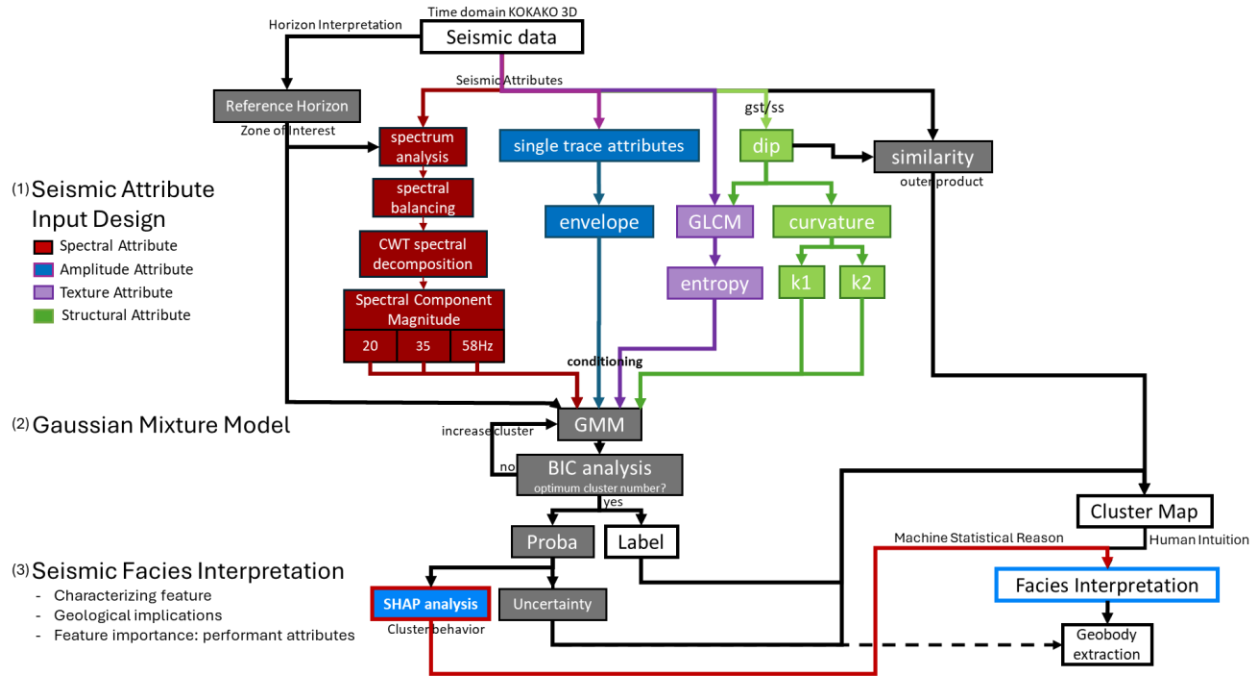


Figure 2.1 Study workflow

The **first component** focuses on constructing a representative, non-redundant input set through careful **attribute selection**. We emphasize avoiding information redundancy between the input attributes to avoid unintentionally introducing bias to the unsupervised clustering process. To do this, attributes are selected by signal expression types—amplitude, spectral, structural, and texture—to ensure that we are capturing conceptually different mechanisms in describing the seismic expressions.

The **second component** applies **GMM unsupervised clustering**, which assigns data points to facies clusters based on normal distributions, allowing for soft classification and uncertainty evaluation. This probability distribution is key for the **third component**, where we introduce **SHapley Additive exPlanations (SHAP)** to quantify the relative contribution of each seismic attribute to the clustering result (probability of membership). This interpretability step enables us to see where clusters differ in their seismic characteristics across the multi-attribute space and make the connection to the implications toward depositional processes, such as channelized flows, submarine fans, or mass-transport deposition that produce these expressions.

2.1 Seismic Attributes

We group the seismic attributes based on their signal expression types: **amplitude**, **spectral**, **structural**, and **texture**. By organizing attributes into these four types, we minimize redundant information or signal variabilities being fed into the unsupervised machine learning. This prevents the model from unintentionally emphasizing or prioritizing certain anomaly-variations in a particular aspect of the signal, at the same time reducing unnecessary computational and interpretation load. This approach also ensures that the inputs maintain representation of each distinct aspect of the seismic response, especially for the expected geological interests in the area.

This allows us to focus our optimization on selection and tuning of the seismic attribute algorithms that best capture the intended signal variation, since, as the saying goes, “*garbage in, garbage out*”, the quality and interpretability of machine learning results fundamentally depend on the quality of the input data, in this case the seismic attribute volumes.

This attribute selection framework is further supported by previous studies, which converge on similar attribute sets when analyzing deepwater mass transport depositional systems using machine learning methods (Zhao, et al., 2016; La Marca, et al., 2022).

No	Attribute Type	Attribute Name	Target signal
1	Amplitude	envelope	Reflection strength (impedance contrast)
2	Spectral	20.5Hz magnitude	low frequency anomalies, thicker strata & structures
3	Spectral	35.5Hz magnitude	medium frequency
4	Spectral	58Hz magnitude	high frequency, thinner strata
5	Structural	$k1$ (most positive curvature)	dome, bulges, <i>frown</i> structures, levees and slumps
6	Structural	$k2$ (most negative curvature)	gulleys, troughs, <i>smiles</i> , channels (negative)
7	Texture	GLCM entropy	disrupted, chaotic, weakly conforming layers

Table 2.2 attributes selected for this study and the signatures they are extracting

Additionally, for the machine learning clustering needs, all input attributes are normalized to prevent bias from attributes with larger values and broader ranges. We apply Z-score normalization as it preserves the shape of each attribute's original distribution. This further emphasizes how unsupervised machine learning clustering relies on the clarity of distinction between signal variations captured by the seismic attributes.

2.1.1 Amplitude Attributes

Amplitude attributes represent acoustic impedance contrasts, changes in density and P-wave velocity, therefore they are often tied to lithology or fluid content. Attributes such as envelope or RMS amplitude provide direct measures of signal strength and can delineate features like bright spots and amplitude terminations. Computation of these attributes are typically direct from the seismic data (raw amplitude wiggle/trace)—independent of other attributes—as they are calculated from instantaneous or windowed single-trace calculations.

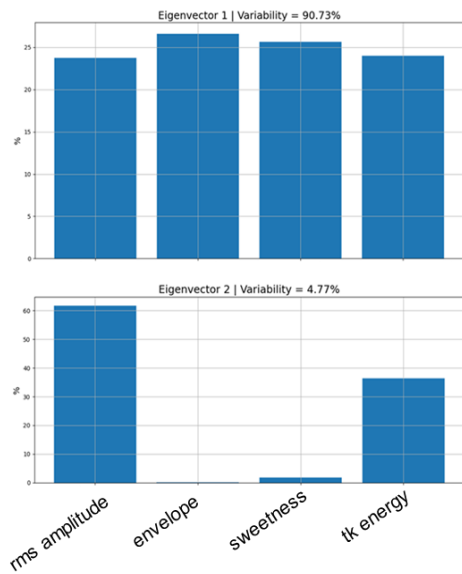


Figure 2.3 Principal Component Analysis eigenvectors variability across various amplitude-type attributes

There is a high cross-correlation among amplitude attributes, confirmed through Principal Component Analysis (PCA) performed on RMS amplitude, Teager-Kaiser energy, envelope, and sweetness; most of the variabilities is captured in the first eigenvector at 90% (Figure 2.3). In PCA, eigenvectors represent the directions of greatest variance in the data, and a dominant first eigenvector indicates that the input attributes largely describe the same underlying variation and carry similar information, reinforcing the importance of this grouping method to avoid redundancy of our input attributes. As such, one attribute is enough to represent this signal-expression type.

Relative impedance, Hilbert transform/quadrature trace, and amplitude volume transform (AVT), although also amplitude-based, are affected by phase domain shift as part of their calculation to infer the layers' characteristics and not necessarily the reflection signal itself. This introduces misalignment in attribute values relative to the reflection events, which the other attributes might not compensate for.

Based on these considerations, the **envelope attribute** was selected, as it represents the complete analysis (both the real and imaginary portion) of the seismic trace and effectively summarizes amplitude-related energy variations in the wavefield.

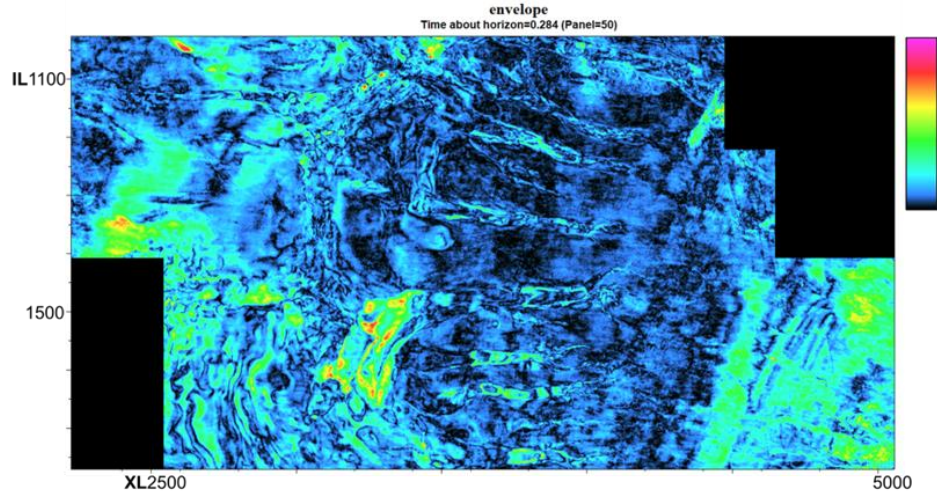


Figure 2.4 envelope (horizon slice 284ms from reference horizon)

2.1.2 Spectral Attributes

Spectral attributes describe the frequency content of the reflection signal in response to the subsurface. These variations are typically due to the interaction of the strata with the moving wavelet that could indicate stratigraphic complexities such as bed thickness (e.g., tuning effects) and/or attenuation due to rock or fluid properties. For example, the magnitude of the 20Hz or 58Hz frequency component can help isolate features with specific vertical-time resolution or layering that was picked up by (interacted with) the longer or shorter wavelength respectively.

We perform spectral balancing to the seismic data (raw amplitude wiggle/traces) prior to generating the spectral attributes to flatten the amplitude spectrum (Chopra, et al., 2016). The resulting attributes represent reflection responses tuned to the subsurface variation corresponding to the spectrum, as such, a flat amplitude spectrum ensures equal energy contribution across all frequencies, effectively broadening the usable bandwidth and enhancing the vertical resolution of the seismic data. The balancing operator filter bounds are adjusted according to the original data spectral distribution on the zone of interest to ensure relevant signal is obtained and not noise; we obtain 10-12-70-80 Hz for the frequency band configuration (Figure 2.5). Then, using Continuous Wavelet Transform (CWT) spectral decomposition, which utilize Morlet wavelets to decompose the seismic volume into bandpass-filtered frequency components (Chopra, et al., 2016), we extract spectral component magnitude banks 7.5Hz wide centered around 13 – 67Hz (Figure 2.5). By feeding the resulting frequency component magnitudes into PCA, we could select the three information-rich frequency

bins at **20.5, 35.5, and 58Hz** (Figure 2.6) to represent the data's spectrum information shown as peaks in variability on the eigenvector variability plot.

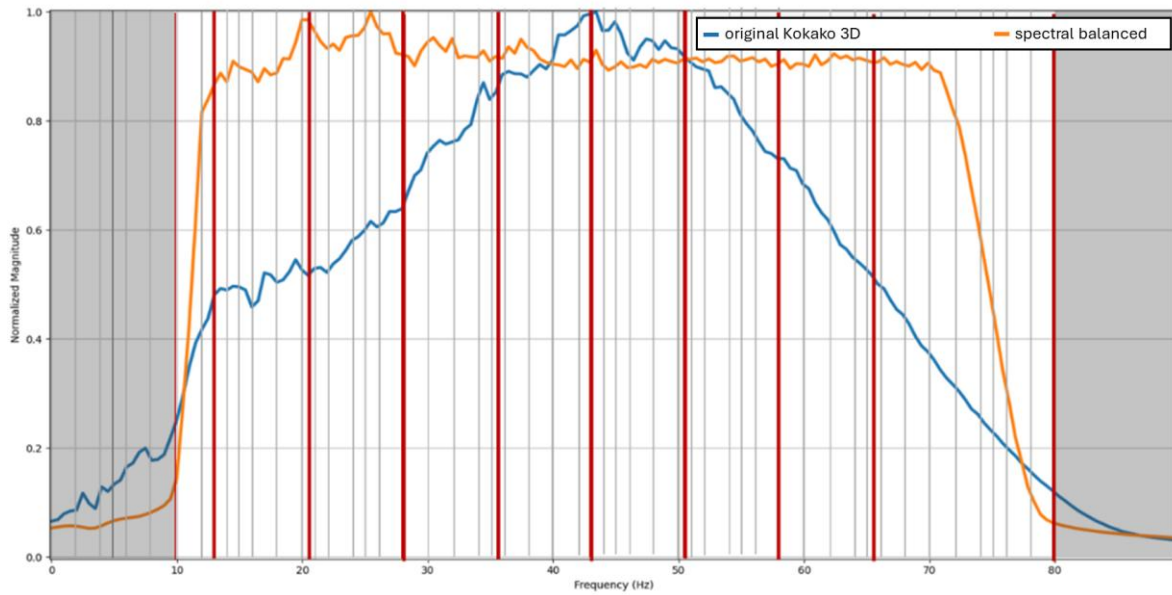


Figure 2.5 Average spectrum from the zone of interest. Blue is the original Kokako3D data, orange is after spectrum balancing is applied. Red lines are boundaries of the frequency bins for the spectral decomposition

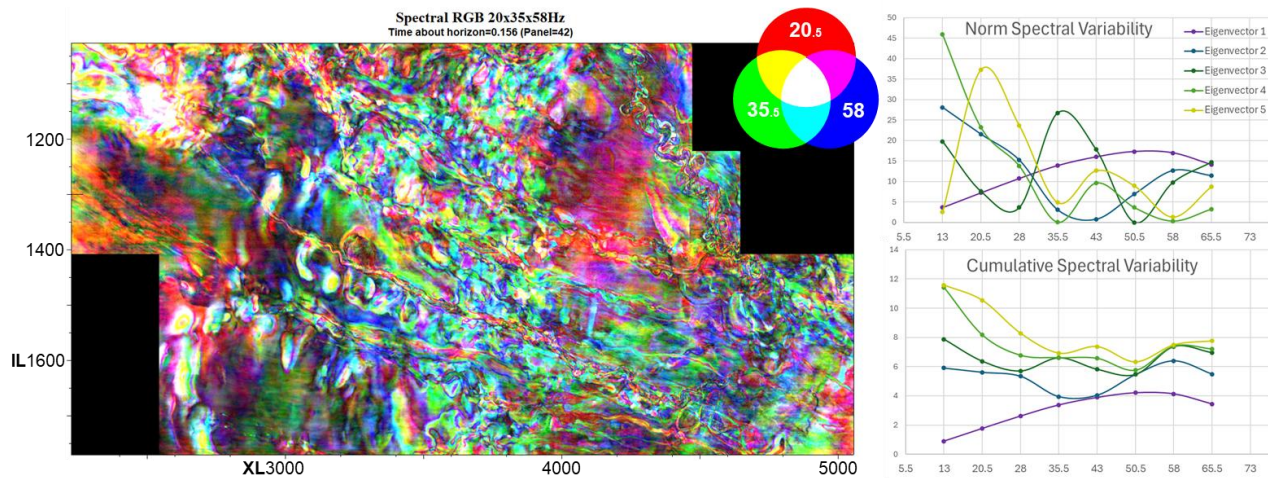


Figure 2.6 RGB Blend map of 20.5, 35.5, 58Hz frequency magnitudes along with the PCA eigenvector values across the frequency bins to show the variability that the three selected bins represent.

2.1.3 Structural Attributes

Structural/Geometry attributes evaluate the spatial organization of reflectors within the seismic volume. Attributes such as curvature or similarity detect features defined not by amplitude but by geometric continuity across traces, revealing elements such as channels, levees, and faults (Roberts, 2001; Marfurt, et al., 1998). These attributes are particularly useful for identifying morphological

expressions of mass transport-related depositional environments, including channel morphology, differential compaction, slump bodies, and overbank deposits.

For the structural and texture attributes calculation, the dip attribute is a prerequisite to guide the orientation necessary for the directional analysis. We used the Gradient Structure Tensor (GST) method to compute the dip attribute (mainly the dip magnitude), as it performs reliably along channel edges with high angular variability (Marfurt, 2006); over the Semblance Search approach, which offers finer lateral resolution but was found to be limited in its algorithm for the large dip angle variations.

Among the suite of structural attributes, ***k1*** (most positive curvature) and ***k2*** (most negative curvature) (Roberts, 2001) were selected for this study to generally represent the reflector flexures and bending, effectively capturing the geometry of channels, slump blocks, and deformation structures.

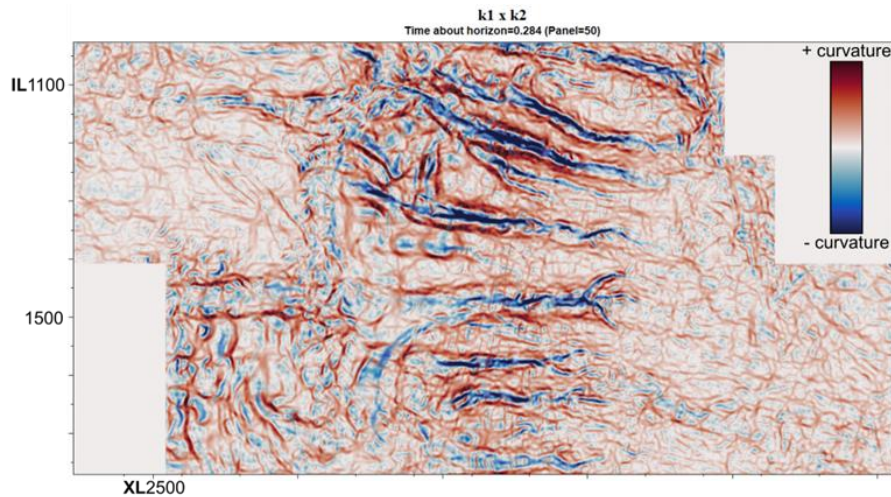


Figure 2.7 *k1 x k2* (corendering) (horizon slice 284ms from reference horizon)

Similarity, or sometimes referred to as coherence attributes, though by definition conform with the structure attribute-type, behave differently in characterizing the seismic data. They measure trace-to-trace waveform continuity, where high values represent laterally consistent reflection events and low values indicate where that continuity is disrupted (Marfurt, et al., 1998). Faults, feature boundaries, and facies terminations/edges are therefore expressed as low-value anomalies, describing the lateral discontinuities rather than characterizing the reflection themselves. As such, similarity serves better as its own independent interpretive dimension—as a co-rendering tool—

useful for delineating visual boundaries of facies distributions on maps rather than as a direct quantitative input for facies characterization and machine learning input.

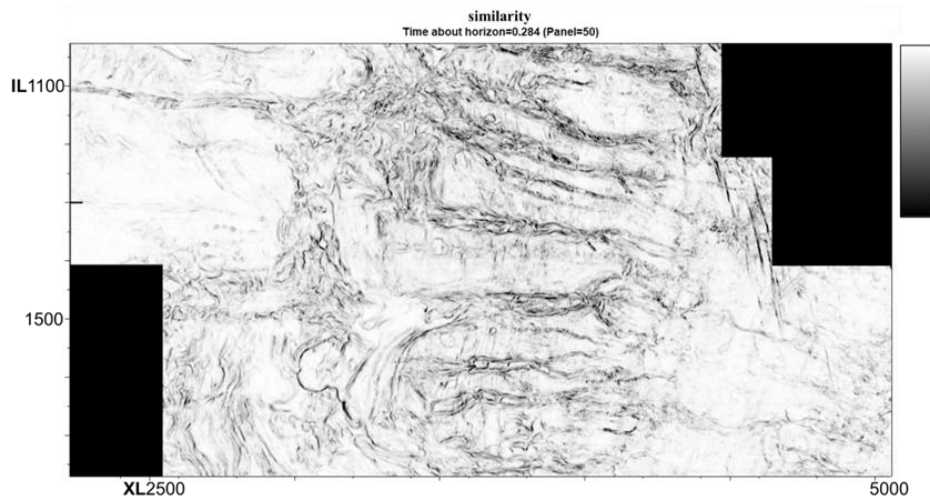


Figure 2.8 outer product similarity (horizon slice 284ms from reference horizon)

2.1.4 Texture Attributes

Texture attributes describe the local variability and patterning within the seismic data. Using algorithms like the Grey Level Co-occurrence Matrix (GLCM) (Haralick, et al., 1973), we quantify how seismic trace wiggles change over a small neighborhood, capturing subtle patterns such as striping or chaotic reflections that may point toward depositional flow regimes, internal structures, or lithological cues.

Haralick et al. (1973) demonstrated the derivation of 14 textural measures from GLCM, each representing different image properties such as coarseness, contrast, and texture complexity. However, it is noted that mainly four measures—*contrast*, *energy*, *entropy*, and *homogeneity*—are sufficient to have discrimination without redundancy (Chopra, et al., 2006).

Several of the derivations offer limited utility for generalized seismic analysis, such as GLCM *mean* which behaves as a low-pass filter. GLCM *contrast*, along with *dissimilarity*, are noted to produce results analogous to the *coherence* family of attributes already part of the workflow. GLCM *variance* is similarly excluded, as it too is closely related, if not identical, to semblance-based *similarity* (AASPI Consortium, 2023).

Among the remaining measures, GLCM *energy* quantifies textural uniformity, whereas GLCM *entropy* behaves inversely, producing large values for texturally non-uniform images, highlighting

areas with high complexity and disorder of seismic reflection. GLCM *homogeneity*, though also useful for quantifying the reflection continuity, measures the overall smoothness of the amplitude differences rather than purely the pattern structure of the signal (Chopra, et al., 2006).

Among these, **GLCM Entropy** was selected for this study to highlight regions of high local signal disorder, which commonly correspond to disrupted bedding, slump deposits, or weakly conforming layers.

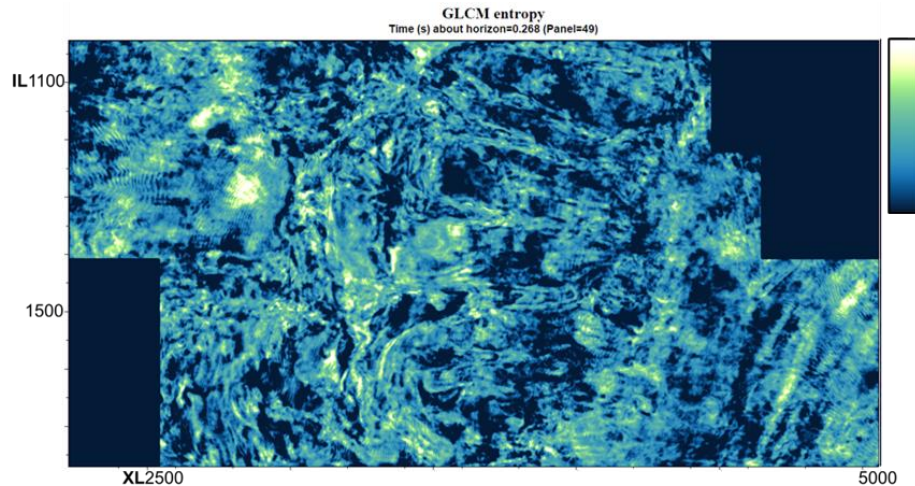


Figure 2.9 GLCM entropy (horizon slice 268ms from reference horizon)

2.2 Gaussian Mixture Models (GMM)

The **Gaussian Mixture Model (GMM)** unsupervised clustering is a machine learning clustering algorithm that fits a combination/mixture of Gaussian distributions (normal distribution) to the spread of data points (McLachlan, et al., 2000). Each Gaussian-cluster represents a distinct statistical population whose shape and position are defined by its mean and covariance.

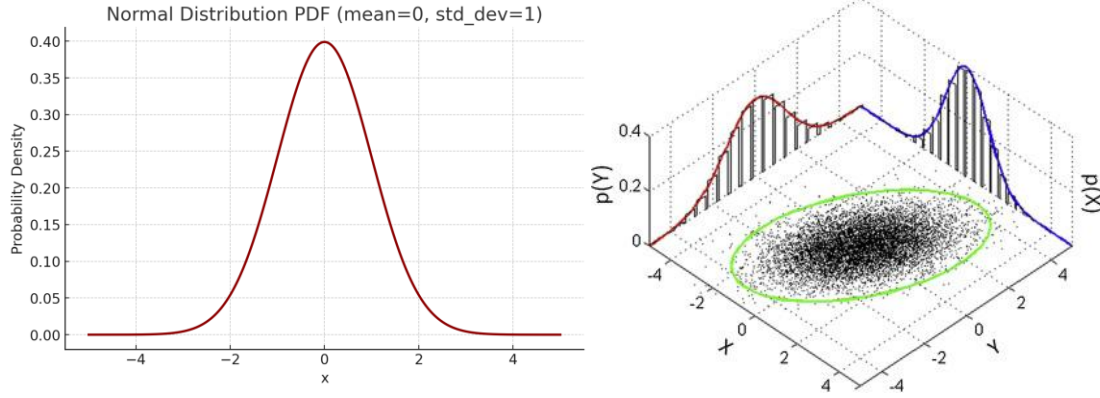


Figure 2.10 An example of a Gaussian normal distribution (probability density function) with a mean value of 0 and standard deviation of 1(left). An example 2D Gaussian distribution of data points with a skewed oval direction indicating a significant off-diagonal element of the covariance matrix(right)

A simple covariance matrix for two dimensions-axis x and y written as follows:

$$\Sigma = \begin{bmatrix} \sigma_x^2 & \sigma_{xy} \\ \sigma_{yx} & \sigma_y^2 \end{bmatrix}$$

Where σ is standard deviation thus the diagonal elements σ^2 is variance—how wide the distribution is in one dimension—and the off-diagonal elements σ_{xy} is covariance which results in the ‘tilt’ seen on the 2D distribution of the data points across the two dimensions.

While each individual facies-cluster may approximate a Gaussian distribution (GMM assumption) the combined distribution is expected to be non-normal, reflecting the superposition of multiple population-groups (McLachlan, et al., 2000) corresponding to distinct lithology/seismic facies (Figure 3.1). This allows for the identification of features like channel structures or hydrocarbon accumulations, which are expected to have attribute values that become a significant variability in the data deviating from the background or other signatures because of their distinct properties.

The parameters of these distributions—mean and covariance—are iteratively optimized using the Expectation-Maximization (EM) algorithm to recreate the original distribution when combined (McLachlan, et al., 2000). This allows for soft clustering of data points by probability density to account for overlapping distributions, making it more suitable for geological data than hard clustering methods, as facies boundaries are rarely sharp and natural distribution of properties often

exhibit smooth transitions (i.e. sand – shale properties) when taking non-uniqueness and heterogeneity into account.

2.2.1 Bayesian Information Criterion

To decide the number of clusters for our GMM, we utilize Bayesian Information Criterion (BIC) to quantitatively find the optimum machine learning model complexity (McLachlan, et al., 2000), in this case, the number of cluster-Gaussians that the model uses. BIC enables evaluation towards how well the model captures the original data distribution while penalizing the number of Gaussians being used. It is calculated through this equation:

$$BIC = k \cdot \ln(n) - 2 \ln(L^{\wedge})$$

$k \cdot \ln(n)$ is the penalty term

where n is the number of data points, and k is the number of parameters in the model (more parameters mean more model complexity). For each Gaussian cluster, k grows by three parameter types: a mean vector (one value per attribute dimension), a full covariance matrix (capturing correlational spread between attributes, contributing $d(d+1)/2$ unique values per Gaussian), and one mixing proportion.

Whereas $2 \cdot \ln(L^{\wedge})$ is the reward term

Where $L^{\wedge}$ is the likelihood, the probability of jointly observing the data points with the model, computed as the product of each datapoint's probability across all Gaussians. Applied to $\ln$ converts this product into a summation of log-probabilities, turning very small probability numbers (<1) into a manageable negative number (log-likelihood). That way it becomes a measure of goodness of fit where the better the fit, the less negative the log-likelihood is. When applied with the negative, this term decreases the BIC number as fit improves.

With this dynamic, the optimal cluster number would then correspond to the point where the BIC value is lowest, or when its rate of decrease converges; it is where an added cluster does not improve the fitting of the overall data distribution enough to justify the added parameter/complexity, if not indicative of overfitting to noise patterns or splitting of existing Gaussians.

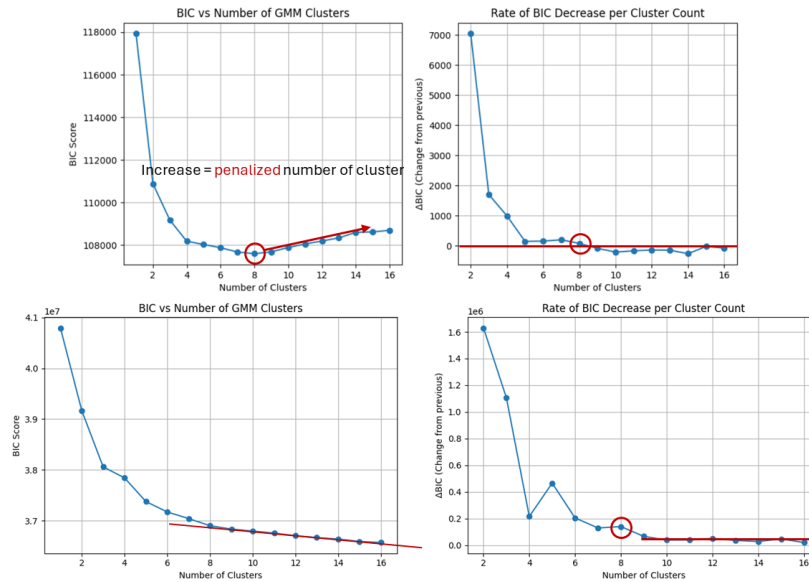


Figure 2.11 How to pick optimum cluster number for GMM using BIC. Top shows a smaller dataset where the BIC value has a lowest point which is at the optimum number of cluster. Bottom is on a larger dataset where the BIC value keeps improving (decreasing) so the optimum cluster number is interpreted at the elbow where the improvement converges. The right graph is the rate of decrease to make the elbow appear more obvious.

It is important to note that the appraised number of clusters does not necessarily represent the count of geological or even geophysical facies of the study area. Instead, it reflects natural statistical patterns to be found in the data distribution. These patterns can take different hierarchical levels of geology—for example, a sinuous channel facies may be represented by its distinct sub-facies such as point bars, levees, or channel banks, while a relatively homogeneous fan system might be represented by only one cluster due to its consistent seismic signature.

Once the optimum number of clusters is found, GMM produces two main outputs: probability values and label. The label volume marks the cluster assignment of the data points (by most probable), which when seen on the volume will show the distribution of the distinguished seismic signals; like what we find on other unsupervised machine learning methods such as SOM, GTM, and k -means. On the other hand, the probability volume gives the degree of membership of a data point towards a cluster, which also provides a measure of uncertainty and transition/mixing between facies. This value will be useful as input of the SHAP analysis to obtain the machine learning cluster ‘behavior’ which allows insight towards the indicative attribute characteristics and ultimately the seismic facies interpretation. It also represents the probable distribution of a desired interpreted subsurface geological object or signal of interest useful for quick extraction and analysis.

2.3 SHapley Additive exPlanations (SHAP)

SHapley Additive exPlanations (SHAP) originates from game theory, where it was developed to fairly distribute rewards among players in a cooperative game (Lundberg, et al., 2017). In this analogy, each seismic attribute can be viewed as a “player” contributing to the final score or probability of “winning” into an output classification by the machine-learning model. The SHAP value measures how much each attribute contributes to the outcome by calculating its marginal effect by testing its presence or lack of presence in the combinations—or *coalitions*—of attributes.

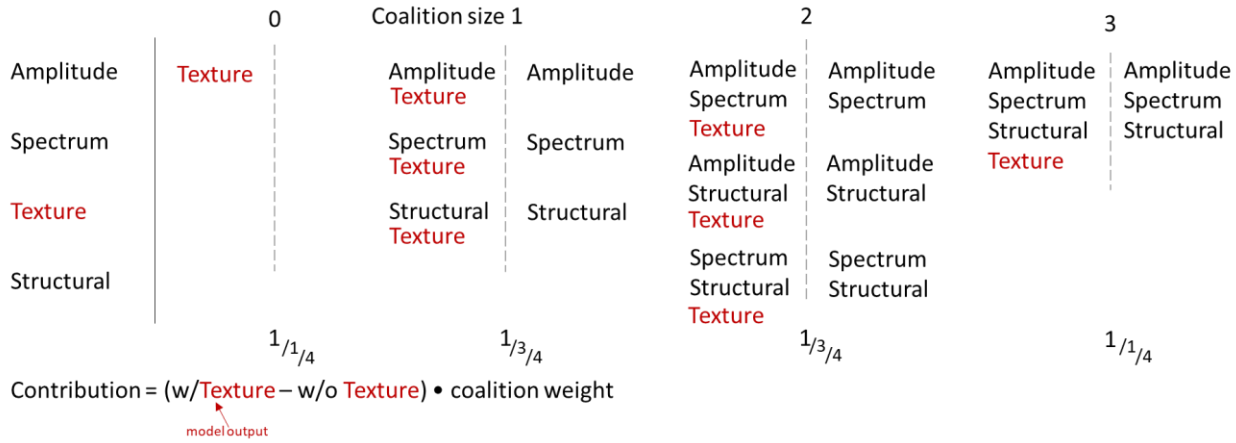


Figure 2.12 How SHAP works by creating coalitions and producing the outputs with and without the feature being analyzed. Then, it weights the result by the number of coalition size (in this example 4) and how many possible coalitions for each size so that each coalition size has the same weight towards the final SHAP value.

Given that SHAP functions by creating a linear model using SHAP values to approximate machine learning model outputs, SHAP has been predominantly applied to supervised machine learning models to uncover relationships of input features towards known output values by a trained machine learning model. In seismic multi-attribute analysis, SHAP has proven to be useful for examining how attribute inputs drive classification toward a trained target label (interpreted facies), uncovering which attribute is key in facies prediction (contributes to the winning probability) (Lubo-Robles, et al., 2020).

In this study, we tweak the logic slightly to where we let the machine learning decide the output labels through unsupervised clustering, then use SHAP to observe how the model's confidence changes in assigning those clusters as the input attribute values vary. By applying the workflow to an unsupervised machine learning clustering method, we achieve a data-driven seismic facies differentiation and appraisal, rather than relying on pre-defined facies labels. This fulfills our aim to

identify the distinct facies that the algorithm can detect within the dataset and utilize SHAP to inspect how the model creates that distinction.

The Gaussian Mixture Models (GMMs) unsupervised clustering method is uniquely suited for this implementation as not only does it produce unsupervised cluster labels, but it also produces a degree of membership towards that cluster in the form of probability density function (Gaussian normal distribution) as part of how the model works. We can thus utilize this cluster probability values as surrogate of the facies membership ‘predictions’ for the SHAP algorithm to observe, where in this case higher probability means a higher confidence in distinction of those facies. This approach enables us to quantify the contribution of each seismic attribute to the probability of a data point belonging to the appraised cluster.

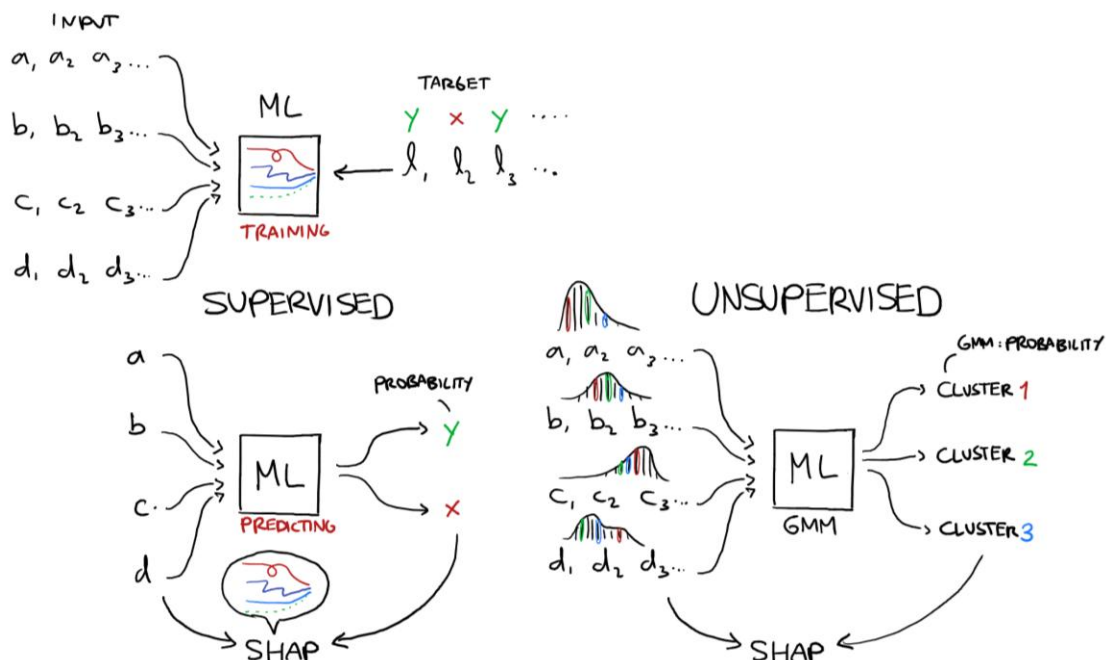


Figure 2.13 How SHAP works and how its logic can be applied to unsupervised ML with GMM.

For example, consider the four attribute types—amplitude, spectral, structural, and texture—working together to determine whether a datapoint is a member of a cluster. SHAP evaluates how the model’s output probability changes when each attribute is added or removed from different combinations of the features — a highly contributing attribute will cause a big change to the output probability whether positively or negatively. The result is a quantity that characterizes how each seismic attribute value causes the probability of membership to increase or decrease for that datapoint toward each cluster. Accumulating this relationship across various attribute values

combinations would then capture how increasing and decreasing of their values pushes into membership or pulls away from membership. This allows a quantitative explanation of how each attribute behaves for the final clustering of data points, effectively characterizing the cluster that was picked up by the GMM unsupervised algorithm.

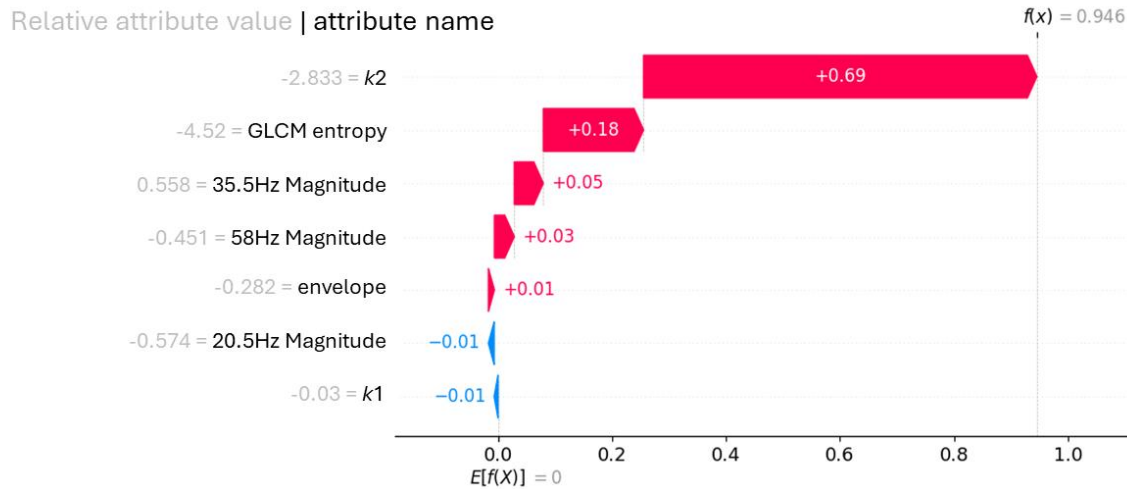


Figure 2.14 SHAP waterfall plot. This graph illustrates what happens behind every datapoint: each attribute input will have a contribution—SHAP value—that adds up towards the final probability value/degree of membership to the cluster. Here we see one of the datapoint that is a member of cluster 4 that is mainly pushed into membership (to the right) by the low relative value k2 (most negative curvature), and low value glcm entropy second.

Of course, directly removing attributes from the model to simulate these coalitions would break the algorithm's structure, since the model uses all input features as context. To address this, the Kernel SHAP approach approximates the SHAP value by sampling a large number of attribute coalitions and substituting the excluded attributes with background values drawn from a reference distribution. The model is then run with these substituted inputs, and the change in output is recorded. By averaging these output changes across many coalition samples, Kernel SHAP estimates each attribute's marginal contribution without ever truly removing it from the model. This method preserves relational structure among attributes, overcomes the missing input problem, and crucially, allows SHAP to be applied to any machine learning workflows such as GMM clustering.

In this study, the background reference distribution is constructed by, first, regularization of the seismic attribute volumes with 16ms, 8 Xlines (CDP), and 8 Inlines increments. Then, randomly sample the dataset stratified by the already-computed GMM cluster output proportions to maintain a representative sample while keeping computation manageable (discussed on appendix A.3). This random sampling was performed for 100 independent batches of 2200 data points each, with SHAP values accumulated across batches to reduce sensitivity to any single random draw and ensure stable

estimates. We used the shap v0.44 python library for the SHAP algorithm implementation, and scikit-learn 1.3.2 python library for the GMM algorithm.

3 Results and Discussion

This section presents the results of the GMM-SHAP clustering workflow and their geological interpretation. The discussion is organized into two parts: first, the GMM clustering results are examined, covering the cluster distributions along with the output cluster map which reveals the geomorphological features—MTDs and turbidites—that hint toward the broader depositional system and allow the clinoform trajectory to be inferred from the interpreted spatial patterns. Second, SHAP analysis is applied to the clustering output to enrich the interpretation of the seismic facies cluster characteristics and evaluate the contribution of individual attributes to the clustering outcome, grounding the facies interpretations in quantifiable attribute relationships beyond intuitive visual recognition alone.

3.1 GMM

Immediately from the attribute histogram overlaid with the clusters' Gaussian distributions, we can see how the GMM algorithm identifies distinct population-groups across the multi-dimensional attribute space. The use of a full covariance matrix for the algorithm allows each attribute to be weighted independently, enabling the model to capture underlying patterns independently across attributes. Each cluster-Gaussian occupies a distinct range of identifying attribute values, and their superposition reconstructs the attribute distribution, this is the defining principle of the mixture model.

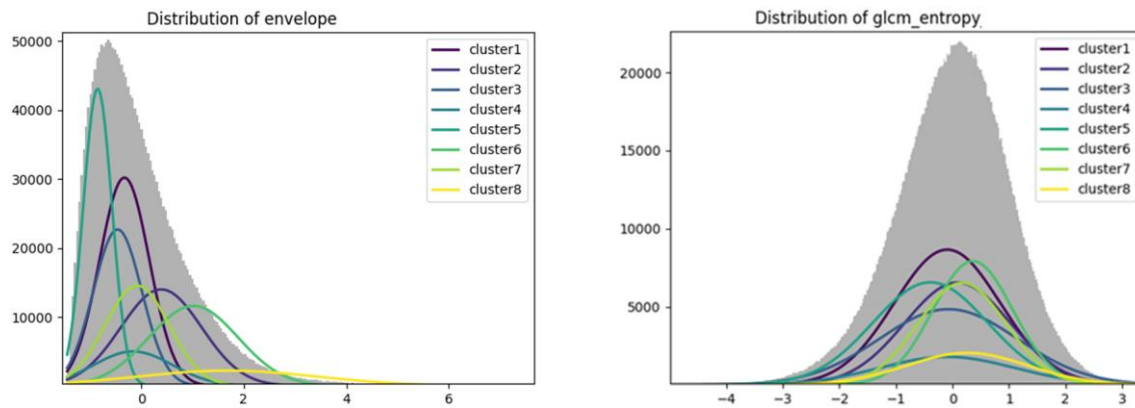
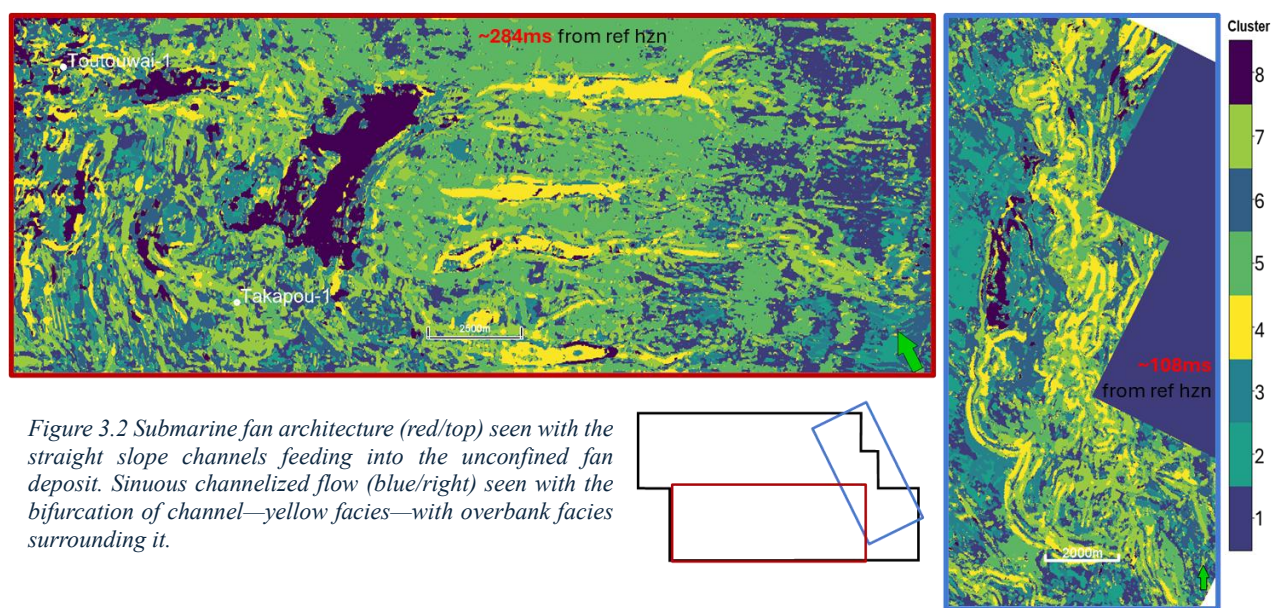


Figure 3.1 envelope (left) and GLCM entropy (right) attribute histogram distribution with the cluster Gaussian distributions overlaid. The combination of cluster Gaussian normal distributions (colored lines) recreate the original non-normal distribution (grey graph), taking into account the juxtaposition of cluster characteristics across attribute values. All histogram graphs can be seen on appendix A.1

The cluster assignments become even more apparent when the data points are viewed on a map as intended. The maps are generated following pseudo-horizon slices based on a reference horizon (Figure 1.3) to follow the clinoform and capture the stratigraphic sequence of events.

This demonstrates how the GMM clusters appraise relevant features useful for geological interpretation. We note that cluster 4 appears to represent the channel facies while cluster 8 expresses submarine fan structures—assigned yellow and dark purple respectively for visual distinction—forming the two primary architectural elements useful for interpretation of the deepwater depositional environment (Figure 3.2).



Two distinct mass-wasting episodes are observable through interpretation of the identified architectural feature clusters, forming separate collections of depositional events that can be spatially mapped to an evolving clinoform breakpoint and slope settings to infer movement of the overall shelf deposition which further extend to infer the shoreline progradation. This is consistent with the two progradational-upbuilding units, LST (Lowstand Systems Tract) and HST (Highstand Systems Tract), where during deposition the breakpoint is most susceptible to instability due to high sediment supply outpacing relative sea level rise and available accommodation space. This instability is most apparent in LST (Figure 3.3) where higher instances of large-scale channels can be seen to initiate along the clinoform breakpoint.

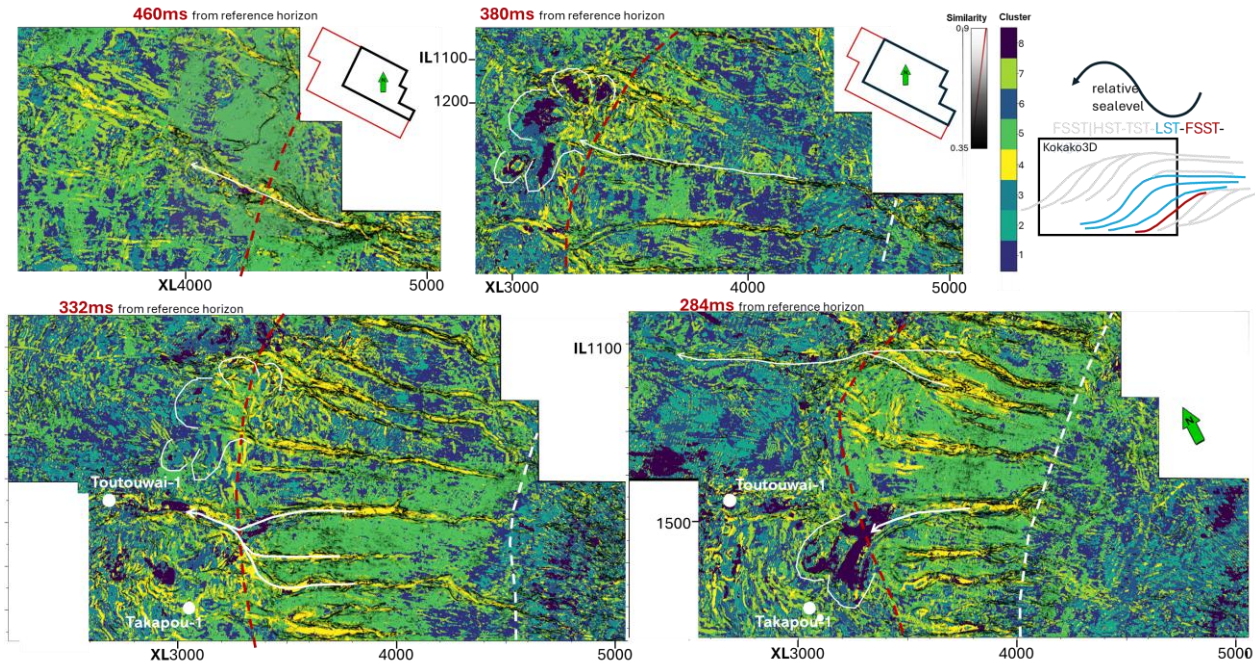


Figure 3.3 The first progradational-upbuilding HST phase where the large straight channels initiates at the clinoform breakpoint and concludes into fan structures at the slope landing. White dashed line show the interpreted clinoform breakpoint and the red dashed line for the interpreted slope landing into the basin floor

It is also observable how the clinoform breakpoint is seen migrating basinward to show not only the progradation but also the steepening of the slope from apparent shortening distance between the interpreted breakpoint line and the basin floor/slope landing. This is characteristic of LST dynamics where, even though still outpaced by the sediment supply, the relative sea level rise is accelerating, causing most of the sediment accumulation to move closer to the breakpoint.

Slice 172ms from the reference horizon (Figure 3.4) appears to mark the end of mass flow events at the basin floor and slope architecture established throughout the majority of the phase. It is followed by a seismic gap without notable features which marks the onset of TST (Transgressive Systems Tract), the retrogradational phase, separating the two upbuilding progradational episodes, where the clinoform breakpoint retreats south-eastward out of the survey area and the survey area observes a uniform deepwater deposition (if any).

The following progradational-upbuilding HST (Highstand Systems Tract) phase sees a more moderate slope, indicated by more sinuous, channelized mass transport occupying what was the topset of the preceding phase, and less occurrences of straight channels broadly following the inherited structural setting of the earlier phase. The buried basin floor morphology of the first phase

likely continues to influence mass flow dynamics and consequently the channel morphology of this later phase.

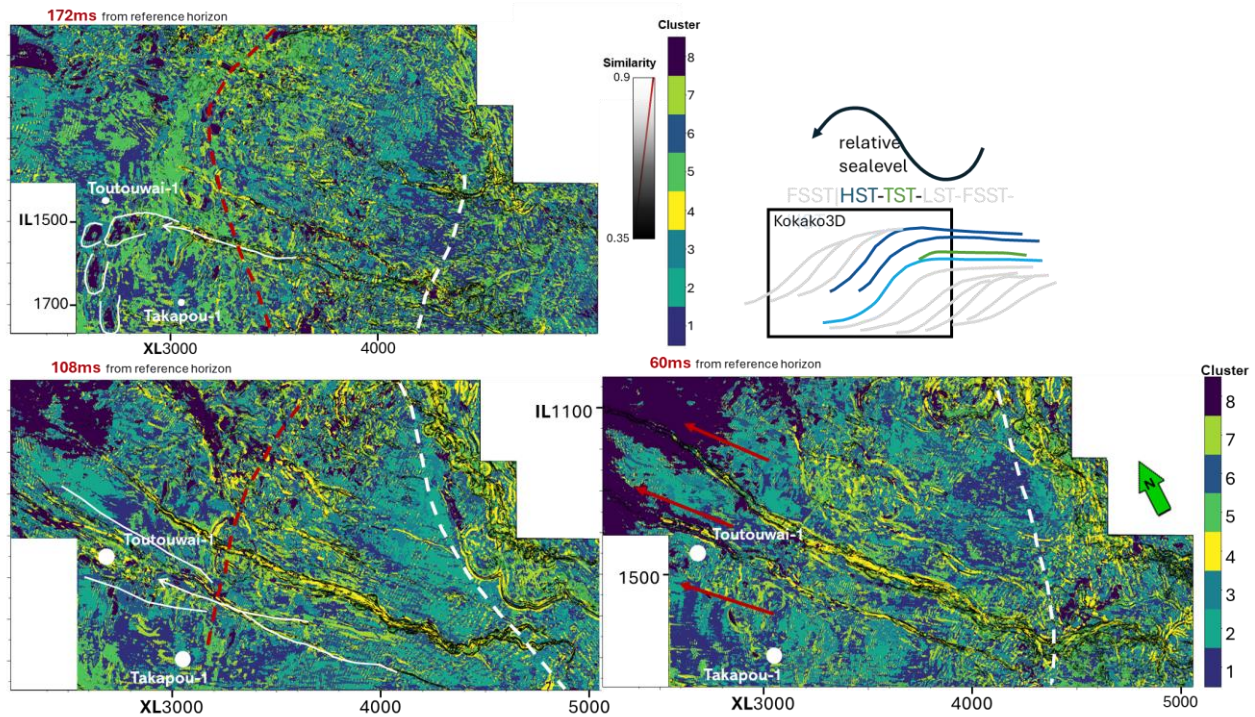


Figure 3.4 The end of the first progradational-upbuilding LST phase (top), later overlain by the seismic gap of the retrogradational TST phase. Then the second progradational-upbuilding HST phase (bottom). Red dashed lines show the slope transition where white dashed lines show the interpreted clinoform breakpoint. The progradation characteristic of this phase can be seen with the shelf landing is not as distinct moving North Westward (red arrows) indicating a gentler slope transition.

Like the previous progradational phase, the general migration of the interpreted slope boundaries is seen to move basinward, although the slope landing appears to become less distinct due to the more moderate steepness of the slope. This is characteristic of the HST phase where the rate of sea level rise is decelerating, and sediment is distributed more evenly throughout the slope.

Eventually the relative sea-level rise inverts, marking the end of HST and the onset of the progradational-downbuilding FSST phase, where accommodation space is being reduced. This transition is marked by the mass transport complex, indicating a shift to a forced regression setting where sediment is driven basinward past the shelf slope, ultimately expressed in the topset architecture through distributary feeder channels and the subsequent erosional surface.

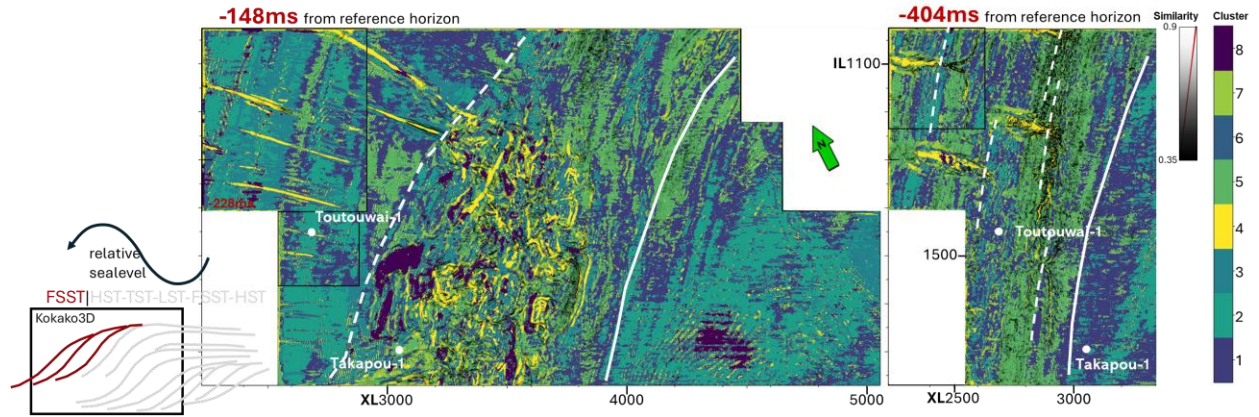


Figure 3.5 the final progradational-downbuilding phase (Giant Foresets Formation) where it's initiated with the mass transport complex and features the distributary feeder channels flowing into the slope gulleys indicating its closeness to shoreline and erosional surface (interpreted with solid white line) from the forced regression

3.2 SHAP Analysis

Having established seismic facies clusters through their spatial expression, we proceed to the SHAP analysis to evaluate whether the machine learning model's assignments are consistent with the expected geophysical responses of the interpreted geological objects. This asserts that the structures identified as 'channel' and 'fan' from the machine learning maps are also as such from the seismic attribute standpoint. In addition, we would also uncover the key signal anomalies defining each structure which could lead toward understanding the geological object properties useful to refer to geological models.

We primarily utilize the beeswarm plot to summarize each cluster's behavior through the calculated SHAP values—the contribution of each attribute to cluster identification. Each dot represents an individual data point's SHAP value plotted on its respective attribute axis, with the population density visualized as the thickening of the swarm. With this, the relationship between the distribution and the relative attribute value—represented by the color—illustrates how the cluster is identified through the attributes across varying conditions in the dataset, creating a graphical description of the cluster according to the machine learning model.

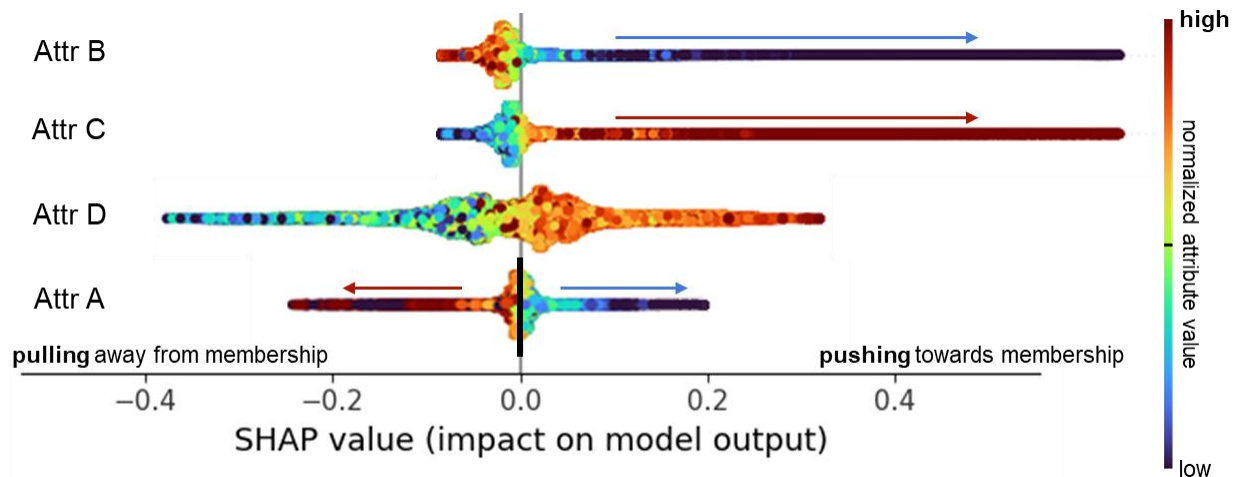


Figure 3.6 Examples of cluster attribute behavior illustrated by global SHAP summary plot: beeswarm plot. The dots represent each datapoint from the observed dataset, its x-axis is the SHAP value of the corresponding attribute, vertically arranged by attribute importance, and the color represents the relative attribute value.

When analyzing the beeswarm plot, there are three things that we can observe about the attribute behavior of the cluster. First, the general direction of the attribute contribution: Observing where the data points are along the horizontal axis—a rightward position indicates a positive contribution pushing towards cluster membership, while a leftward position indicates a negative one pulling away from membership—and cross-referencing with the relative attribute value—represented by the color of the dots—allows us to conclude whether high or low attribute values are supporting or suppressing membership. This directly characterizes the cluster; for instance, Attribute B in Figure 3.6 has low attribute values (blue) plotting to the right, indicating that the cluster is characterized by low values of that attribute, while conversely, Attribute C shows its high (red) values being the cluster characteristic (plots to the right).

Second, polarization of the attribute characteristic: whether the swarm distributes in one direction or both reveals the nature of the attribute's relationship with cluster membership. A one directional distribution, like the one we observe on the example Attribute B and C, means a consistent characteristic across the value range, functioning as an indicator where past a certain attribute value membership is consistently supported or suppressed. Meanwhile, a bidirectional distribution suggests a more complex relationship where the degree of polarization is also informative. For instance, Attribute D in Figure 3.6, shows a relatively even spread in both directions, suggesting a gradational relationship where the contribution scales with attribute value. Attribute A, by contrast, is more polarized, with extreme values contributing strongly to both directions while the general population clusters around the center, indicating a transitional phase where the attribute contributes

minimally to cluster identification. A third scenario arises when the intermediate attribute values are the primary indicators, producing a varied color distribution where the swarm initially shifts in one direction before reversing as attribute values increase or decrease (Figure 3.7); suggesting that a specific range of intermediate attribute values is characteristic of cluster membership rather than the extremes.

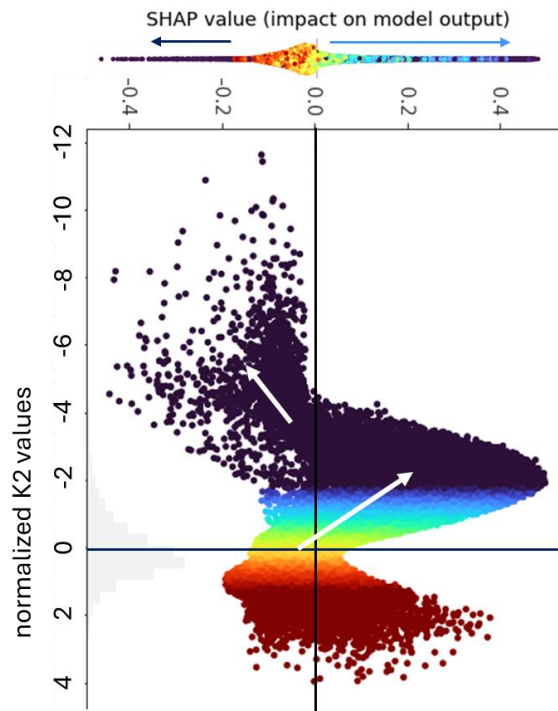


Figure 3.7 k2 (most negative curvature) contribution towards cluster 7 plotted against its normalized value shows the complex behavior of the attribute characteristic in identifying this cluster. We find a generally supporting low value (blue, pushing towards membership—to the right), but the direction changes when moving into the extreme low attribute value. The resulting SHAP beeswarm plot is shown at the top.

Finally, attribute importance shown as the vertical ordering of attributes on the plot, ranked by their mean absolute SHAP value. Attributes ranked higher have more overall contribution to cluster identification regardless of direction, which could also be evident in the wider spread of their swarm. This indicates that the attribute characterizes the cluster well, often meaning that the cluster appears most anomalous from the perspective of that attribute. It does not necessarily mean that the lowest ranked attribute fails to describe the cluster, but rather the distinction is not as contrasted in certain attribute dimensions, hinting at either an underperforming attribute input or an inherently ambiguous property of that cluster.

3.2.1 Cluster Characteristics

With that in mind, we now observe the clusters that were most useful in the depositional environment interpretation and see how the clusters were appraised by the machine learning algorithm.

The first one, **Cluster 4**, seems to represent the channel facies when looking at the general map interpretation (colored yellow on Figure 3.2).

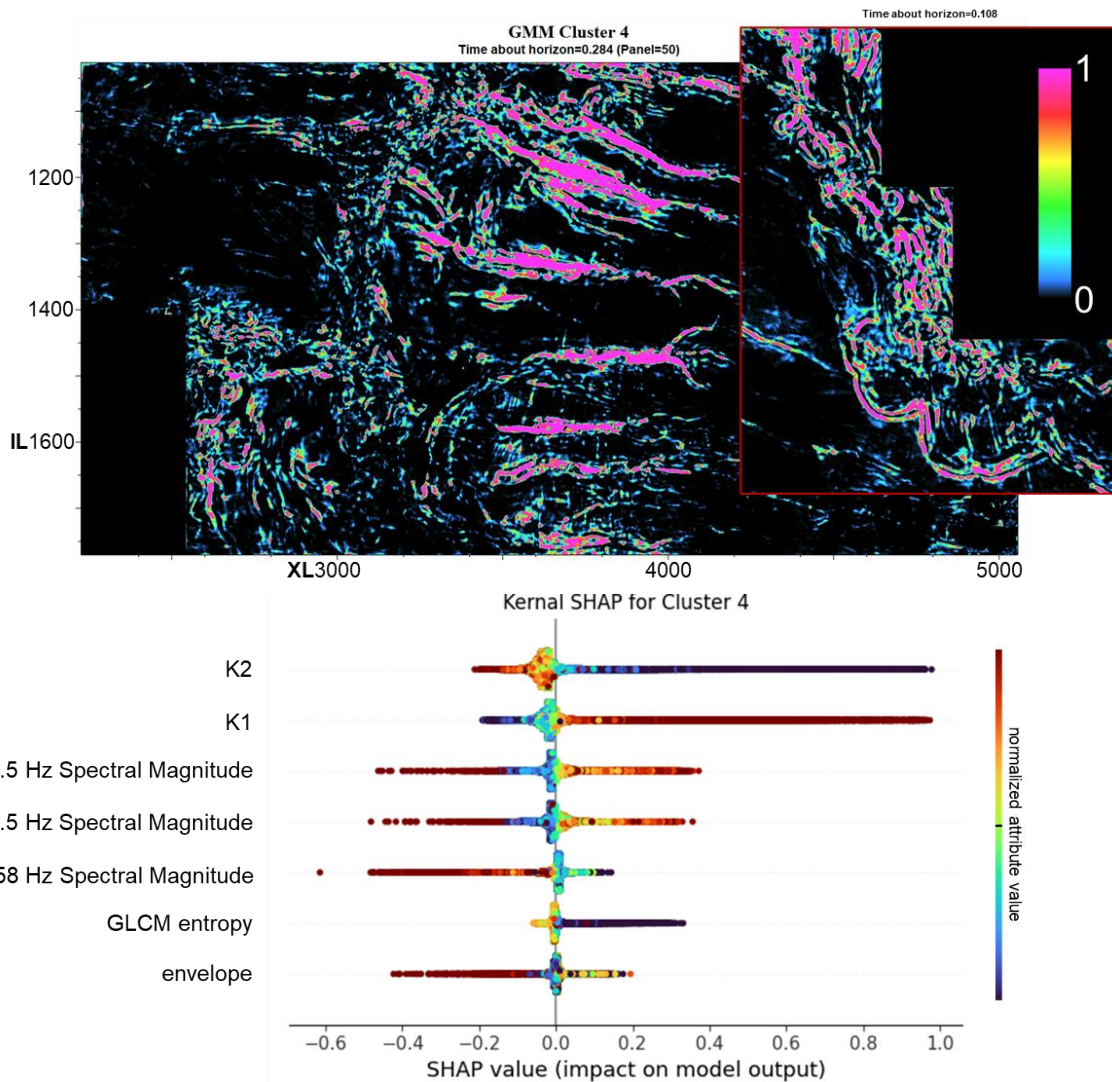
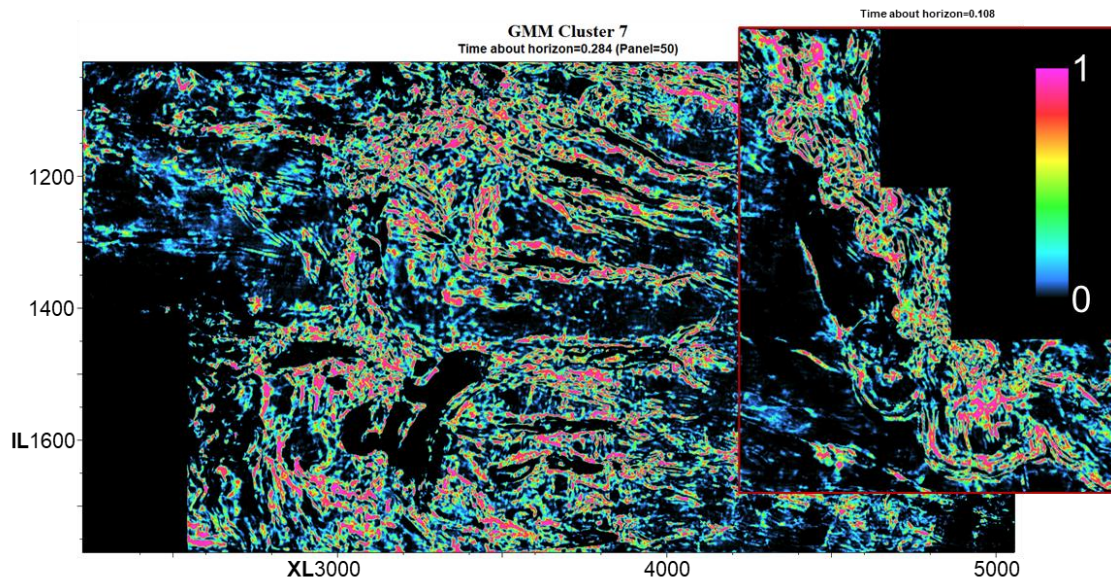


Figure 3.8 Cluster 4: Channel (low entropy) body, mainly characterized by the geometrical input attributes k1 and k2. The high magnitude of low frequencies (20 and 35Hz) supporting the membership while high 58Hz magnitude pulls away from membership further indicate the tuning of vertical layering of the channel structure to the lower frequency bins. Another interesting note is GLCM entropy can be seen supporting with its low values, which means this facies have a well layered, undisrupted, reflection.

According to the SHAP plot, the cluster's most defining characteristic is low value $k2$ indicating a concave or *smile* structure, but interestingly members may also be *pushed* by high $k1$ value, indicating a convex or frown structure characteristic; with almost identical beeswarm distribution and a strong one directional, indicating behavior. This hints that this cluster is mainly driven by a distinct curvature or geometrical characteristic that in this exercise is primarily tuned for what we expect for channel geometry. Levees would produce a positive curvature, and the channel body/trough would appear as negative curvature.

Another point to observe is how the lower frequency attributes (high magnitude of 20.5 and 35.5Hz) is *pushing* towards membership while high frequency—represented by high 58Hz magnitude—*pulls* away from membership. This means we expect the channel structures to have certain deeper/thicker layering that tunes to the lower frequencies. Finally, we also see how GLCM entropy contributes positively when it is of low values, signaling that this cluster also has a well layered characteristic.

Another channel related facies is **Cluster 7**. It is also primarily indicated by the expected channel geometry ($k1$ and $k2$ curvature) according to its SHAP plot, with $k1$ being the slightly more influential attribute, as opposed to $k2$ on Cluster 4.



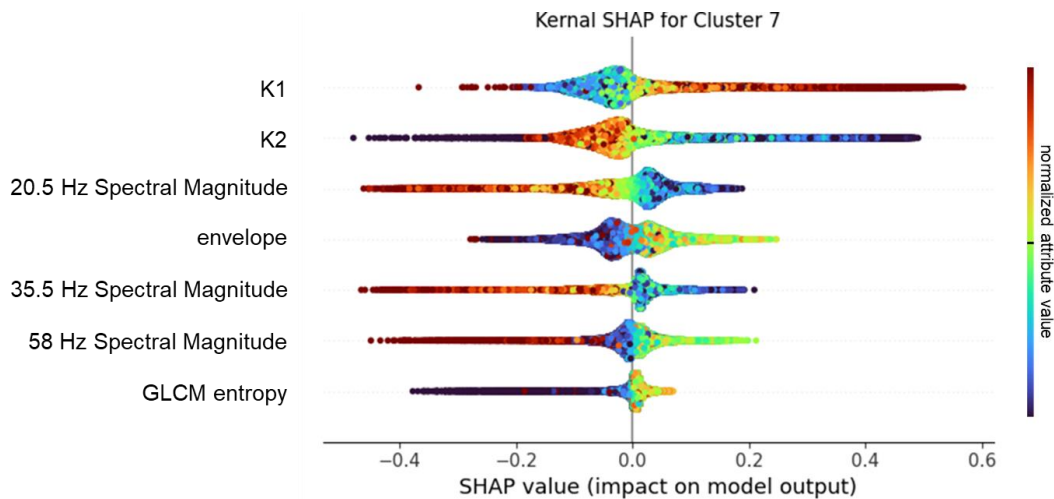


Figure 3.9 Cluster 7: Channel overbank or edges (thin layers). While having also K1 and K2 as the main indication similar to Cluster 4, high magnitude of the low frequencies are negatively supporting instead (pulling away from membership) while intermediate high magnitude of high frequency (58Hz) is pushing towards membership, indicating a tighter reflection layering.

We note that the swarms are generally more spread out than Cluster 4 and some extreme values of $k1$ and $k2$ are seen not supporting the cluster membership, evident in some of the darker hue dots plotted towards the left opposing the general relative value's direction.

We also observe spectral magnitude behavior where high magnitude values of 35.5 Hz and 20.5 Hz pull away from membership, while the low values, along with the intermediate-high 58Hz (high frequency) magnitude pushes towards membership. This indicates that the cluster expects expression of relatively thinner strata with a tighter alternation of lithology contrasts, not thicker ones that tunes to lower frequencies. Hence why it appears that when it finds high magnitude of lower frequencies it negatively contributes away from the cluster.

The last main indication is with the low GLCM entropy pulling away from membership, which could mean a facies signal that is not well ordered. Perhaps a chaotic reflection signal or a complex lithological setting (from thin beds) that causes a disrupted signal. This is further supported by less certainty (probability values) at the edges of the facies body on the map, hinting at a facies that has an underlying complex seismic response.

We look towards Cluster 8 that seems to point at the submarine fan structures that is also key to interpreting the depositional environment. We can see how, especially on the deeper depositional system (at 284ms), it extracted relatively wider structures that spread out of the channel structures.

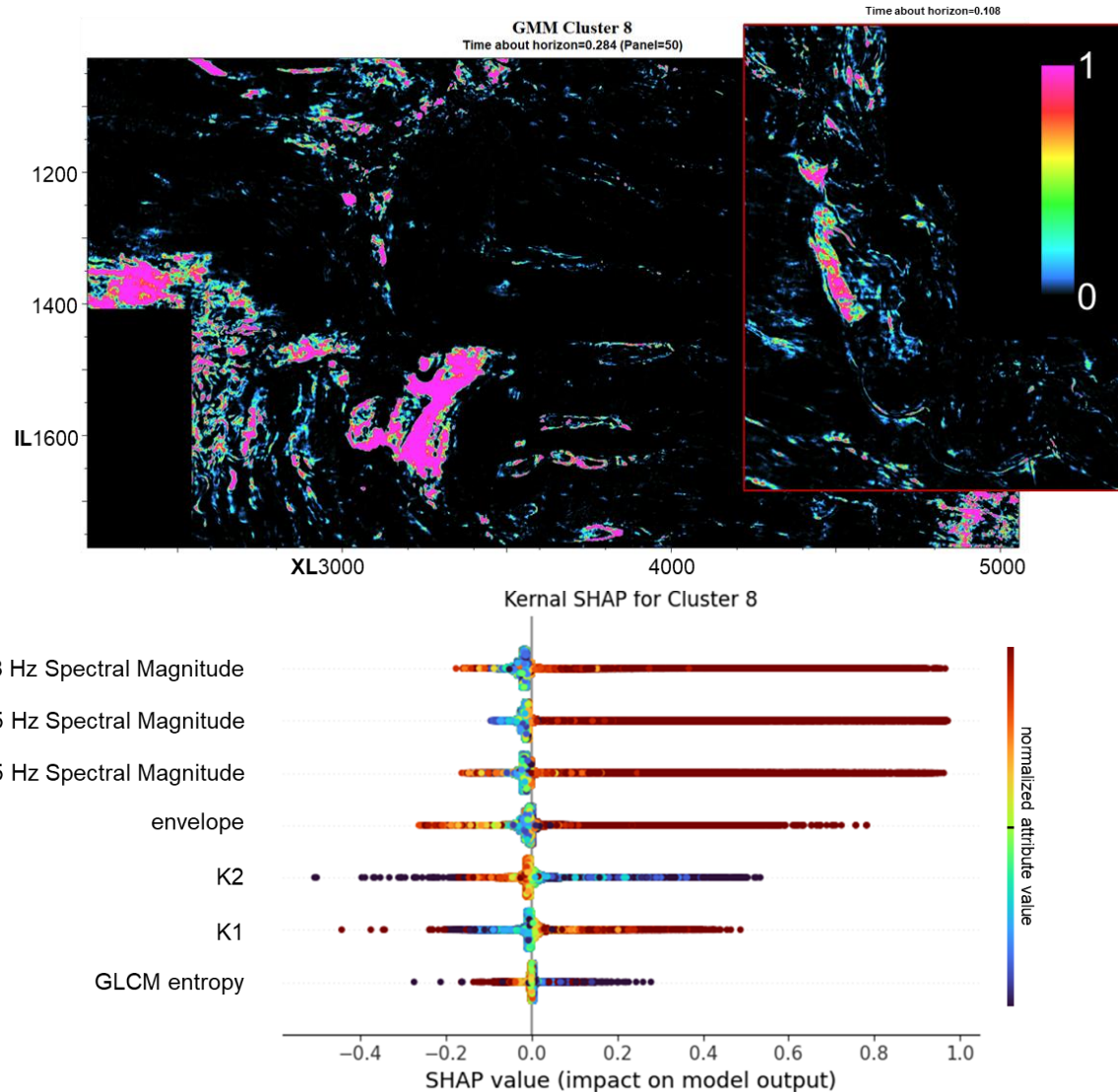


Figure 3.10 Cluster 8: Fan facies, mainly indicated by high reflection energy across the spectrum. GLCM entropy hinting a positive (pushing) contribution from the low values meaning a generally coherent, unbroken reflection spatially. As expected of a submarine fan deposit.

Looking at the SHAP plot, this fan facies is indicated by high spectral magnitude values across the frequencies used. Furthermore, right next on the attribute importance is envelope attribute that also characterizes amplitude of the signal. This strongly hint at this cluster mainly being characterized by high signal energy or strong reflections, indicating a high impedance contrast which could be from a significant lithological change from the overlying layer, as expected of a submarine fan deposit.

Interestingly, we note that it's only the especially high values (dark red) on envelope attribute that positively contribute to membership, while the general population is very closely plotted around zero

and the moderately high is, instead, pulling away from membership. Similar behavior is seen on 58 and 35.5Hz magnitudes.

k_1 and k_2 also exhibit nuanced characteristics, though not as strong when comparing its SHAP value of 0.6 here to Cluster 4 which almost reaches up to 1 (very strong influence).

GLCM entropy, although lowest on the strength of contribution, exhibits a well polarized, two-directional behavior where low entropy pushes towards membership and high entropy values pulls away from membership. This is indicative for a facies that tend to have a persistent and coherent reflection signal.

We can also observe that the clusters coincide with evidence of sand (low GR) through the well logs. On Xline 2693 section (Figure 3.11), where Toutouwai-1 is located, we can see the well goes through the key architectural feature clusters of the first progradational-upbuilding HST phase at about 284ms (Figure 3.3), and the second progradational-upbuilding LST phase at 108ms from reference horizon (Figure 3.4)

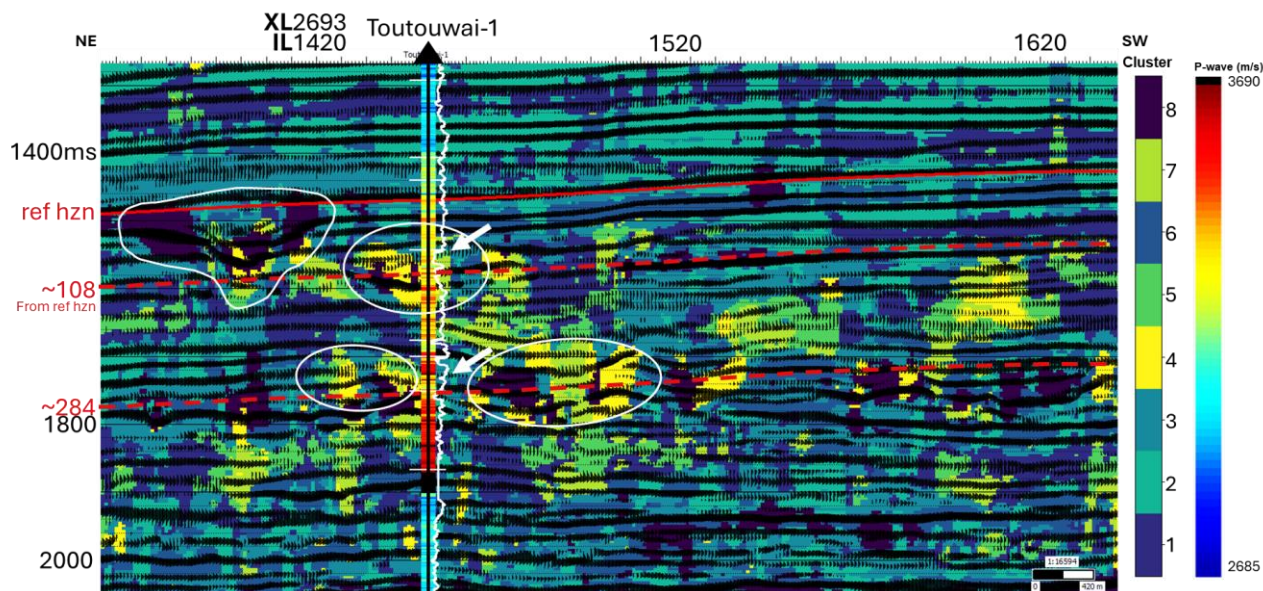


Figure 3.11 Xline 2693 of the GMM cluster output corendered with the seismic trace (peaks), showing the location of well Toutouwai-1. We can see how the well penetrates through some of the key architectural feature clusters. 108ms from reference horizon (top red dashed line), on the well location, the top white arrow points at where the GR log (white wiggle) is low (dips left) which coincides with the channel facies (cluster 4) even when P-wave (colored well column) doesn't seem to exhibit any clear distinction. The bottom white arrow points at another anomaly this time indicated by a low velocity (green-yellow) and high GR which happens to coincide with the overbank facies (cluster 7) on the side of a channel facies (cluster 4).

On the cross section views, we can see how the clusters tend to mark the edges of structures instead of the body itself. This is especially so for the channel facies (cluster 4) where most of the

time the cluster is identified on the bottom of the channel structure, while it struggles to pick up the top of the channel incision through the layers. This is to be expected since the attribute inputs of the machine learning model are of the reflection signal—layer contrasts—thus it characterizes the seismic facies boundaries instead.

Takapou-1 can be seen on Xline 3057 (Figure 3.12) where it seems to miss our main architecture element clusters (cluster 4 and 8) except for one channel-like structure at about -148ms from the reference horizon, where on the slice (Figure 3.5) it seems to be a portion of the mass transport complex at the onset of the downbuilding FSST phase, the end of the progradational-upbuilding HST phase.

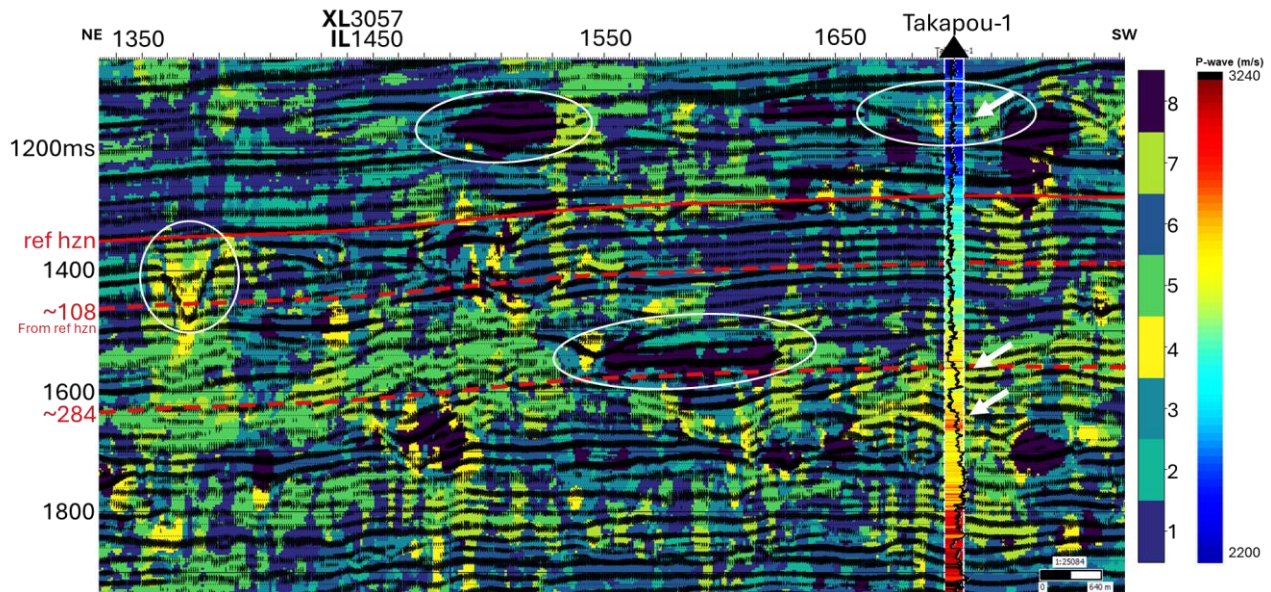


Figure 3.12 Xline 3057 of the GMM cluster output corendered with the seismic trace (peaks), showing the location of well Takapou-1. The large fan lobe (cluster 8) at 284ms from reference horizon (the lower red dashed line) can be seen just to the left (NE) of the well location, while the well itself seems to go through the complex mix of facies with GR highs and lows (black wiggle inside the well column) and velocity anomalies (marked with the two lower arrows) indicating a complex layering of different lithologies. The top arrow points at the mass transport complex channel-like structure at about -148ms coincide with a slightly high p-wave velocity on the well log and marked as channel facies (cluster 4).

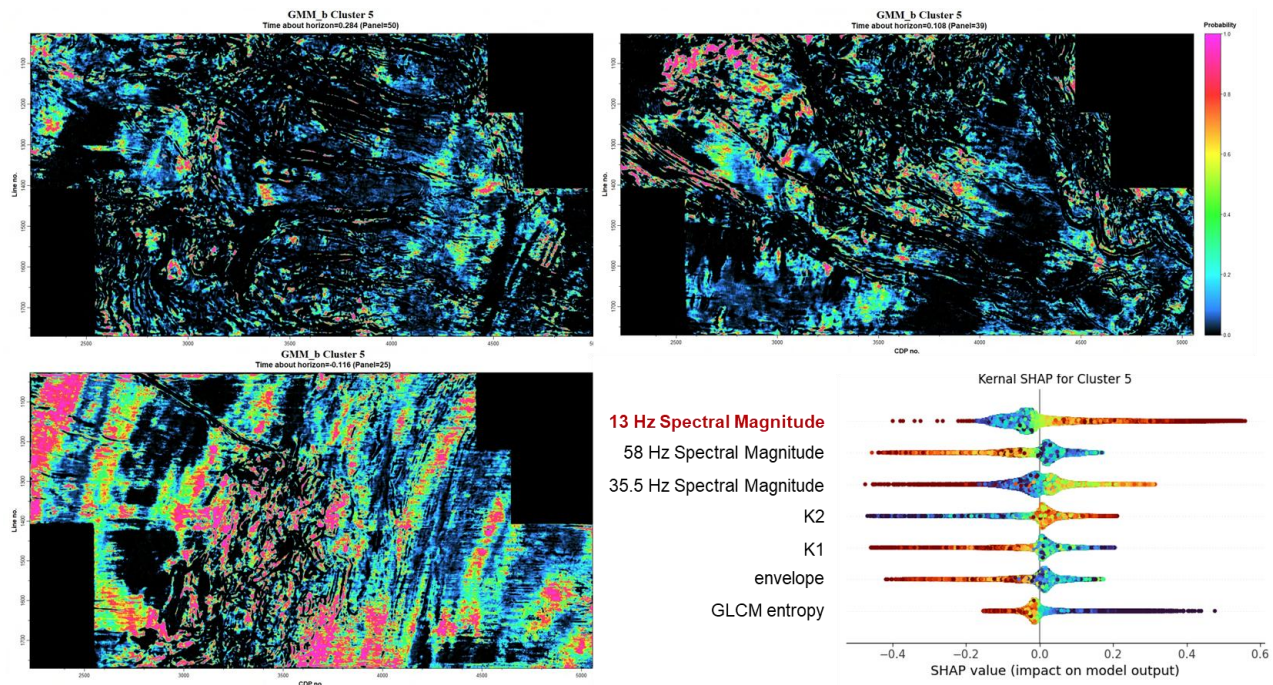
With this we can see how machine learning's multi-dimensional analysis was able to successfully pick up the architectural elements, making well-log and vertical section interpretations easier to connect with the larger seismic spatial context. It leverages the combination of geometrical attributes when velocity anomalies, consequently amplitude related attributes, were ambiguous and vice versa; confirmed by what the SHAP analysis described of each cluster characteristics.

3.2.2 Attribute Contribution Performance

SHAP analysis could also offer a way to evaluate how well a selected input attribute performs in the unsupervised clustering.

Like how a football coach can tweak a player's training for certain abilities according to their performance and roles in the playbook, with information from the SHAP analysis we can tweak the input attributes parameters to produce better—more defining—contributions to the facies cluster, improving extraction through better signal characteristic distinction.

We demonstrate this with one of the spectral magnitude attribute inputs. When utilizing the 13Hz spectral magnitude as one of the attribute inputs, the GMM workflow produced the cluster 5 as follows:



We find that, for this cluster, the 13 Hz spectral magnitude attribute dominates as the highest-ranked attribute by importance yet contributes minimally and behaves inconclusively across the other clusters. Within this cluster specifically, it is the only attribute contributing positively toward membership, while the remaining attributes contribute negatively — suggesting that this cluster functions as a residual grouping, where the other attribute indications support membership into other clusters instead. When projected onto the map, the cluster exhibits a nondescript signal distribution,

frequently manifesting as reflection bands coherent with the 13 Hz waveform. This suggests that the attribute is describing its own pattern rather than a meaningful geological signal, effectively capturing background noise coherent to the 13 Hz banding rather than a genuine facies expression.

A similar evaluation could also be applied to the structural attributes—most positive ($k1$) and most negative ($k2$) curvature attributes or specific shape indexes—which are tuned to extract specific architectural geometries. These attributes are expected to contribute meaningfully to the differentiation of structurally distinct facies; their contribution would therefore reflect the degree to which the target architectures (signals) are expressed or enhanced at the current attribute parameterization. This logic extends to any input seismic attribute type, provided its parameterization can be tuned to enhance specific geological expressions.

This illustrates a broader utility of the SHAP-GMM framework beyond geological interpretation: as a diagnostic tool for iterative attribute selection and parameterization, guiding refinement of the input attribute set toward more geologically targeted and discriminative inputs.

4 Conclusion

This study demonstrates that data-driven seismic attribute characterization can be obtained from unsupervised machine learning on the Kokako 3D Seismic data of the deepwater depositional systems of Taranaki Basin, New Zealand. This is done by integrating SHapley Additive exPlanation (SHAP) into Gaussian Mixture Model (GMM) clustering of a multi-attribute analysis workflow. Global SHAP summary plots serve as the main tool that accumulate SHAP values across multiple data points to communicate how the defining seismic attribute values contributes to the identification of the output machine learning clusters. This bridges the interpretation of GMM clustering results with the seismic attributes that constitute its input, ultimately linking the signal characteristics measured by the seismic attributes to the algorithmically appraised clusters, allowing seismic attribute theory to inform geological implications.

As such, careful deliberation of input attributes is key to the performance of unsupervised machine learning clustering in capturing geological signatures of interest. The quality of contrast of the variabilities that the input seismic attributes can render directly translates to how effectively they contribute to cluster identification. For the same reason, the seismic attribute characterization derived from the SHAP summary plots is also governed by the combination of input attributes selected. The seismic attributes used in this exercise measure and characterize reflection interfaces rather than the layers themselves, and this distinction should inform both the interpretation of results and the attribute selection in any future application of this workflow.

Taking that into account, the cluster characterizations by the SHAP method align with the expected seismic attribute responses of the interpreted geological objects of the clusters. This suggests that architectural elements of the Taranaki deepwater depositional environment extracted through GMM unsupervised clustering are reliable enough for constructing interpretations of the subsurface history. By identifying the attributes most critical to each seismic facies, SHAP provides a quantitative basis for multi-attribute seismic interpretation. Additionally, it offers an intuition-driven mechanism for quality control and trust in machine learning products, especially for unsupervised methods without prior training or labeled guidance.

4.1 Limitations

While the results demonstrate the viability of the SHAP-GMM workflow for seismic facies characterization, several limitations inherent to the data, methodology, and scope of this study should be acknowledged. First, in the interest of algorithmic simplicity and controlled testing, this study utilizes a single post-stack 3D seismic volume, avoiding complications arising from spectral matching, non-zero incidence angle heterogeneity, and migration inconsistencies that come with integrating multiple volumes as with prestack or angle stack cases. As an extension of this simplicity constraint, the study focuses on reflection-based attributes that characterize impedance contrasts directly, without prior model assumptions or inversion processes. Consequently, the results are oriented toward evaluating reflection interfaces rather than the intrinsic properties of the geological layers themselves.

A further limitation arises from the nature of GMM unsupervised clustering, where the algorithm is free to assign arbitrary labels to the facies it identifies. Proximity in cluster numbering does not imply similarity in characteristics as the clustering algorithm does not preserve the topography of the underlying attribute space. Variations in the dataset that the model observes—data point extraction, preconditioning applied, random states—would affect the clustering outcome. As such, comparison of clustering outcomes across different run configurations would require careful re-examination.

4.2 Recommendations

The findings of this research of unsupervised clustering SHAP open several possibilities for future study:

First, the current exercise exclusively utilizes seismic attributes that analyze reflections, characterizing interfaces rather than the layers themselves. As consequence, so does the machine learning output. Incorporating layer-evaluating attributes, such as those derived from seismic inversion or mathematically transformed to be attributed to the indicated stratigraphic layers, could guide cluster characterizations more toward the geobody and broaden the geological implications that can be drawn from the SHAP analysis.

On that note, incorporating amplitude versus offset (AVO) or elastic impedance attributes would also add further dimensions of characterization. Where prestack or angle stack data are available, this

information would contribute independent signal variations related to fluid content and lithological heterogeneity; facies distinctions that post-stack attributes might not be able to capture.

Second, the unsupervised clustering SHAP integration in this study is primarily enabled by GMM's unique Gaussian membership probability output. Extending this approach to other unsupervised clustering methods would require an analogous continuous membership or a distance metric to serve as an equivalent SHAP target. The development of which would generalize the framework across the broader unsupervised machine learning toolkit available to seismic interpreters.

Finally, the SHAP summary plots generated in this workflow present a structured and information-rich output that could serve as a base for automated interpretation. A foundation model trained on SHAP graphic outputs could then learn to recognize characteristic attribute contribution signatures associated with known geological facies, rather than directly training models to recognize patterns on raw seismic data.

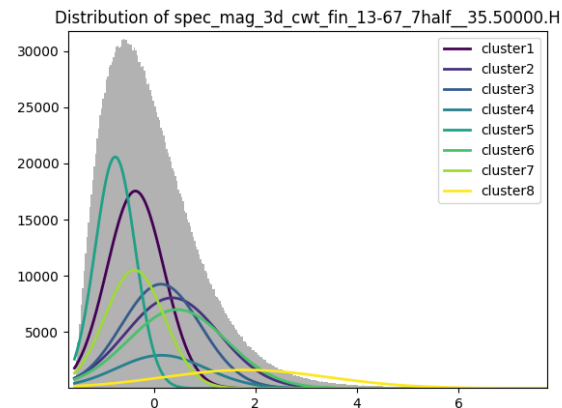
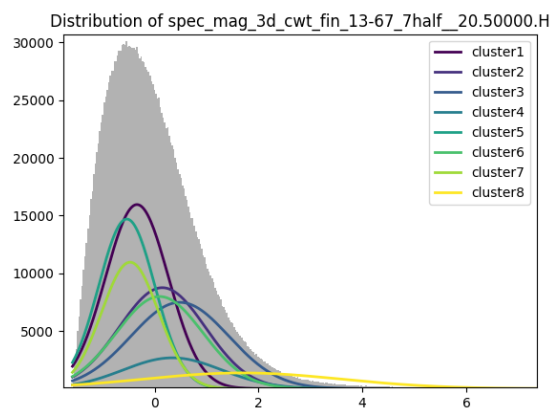
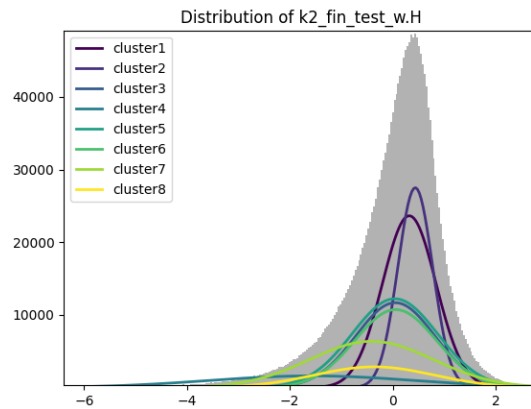
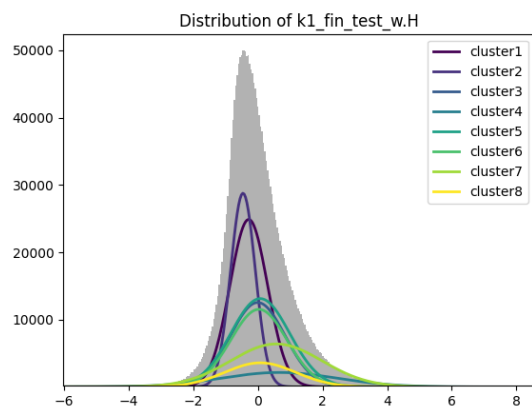
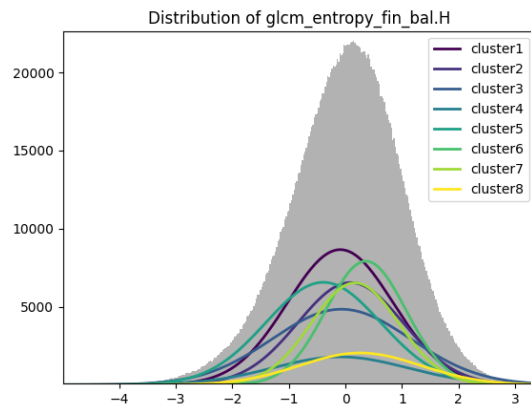
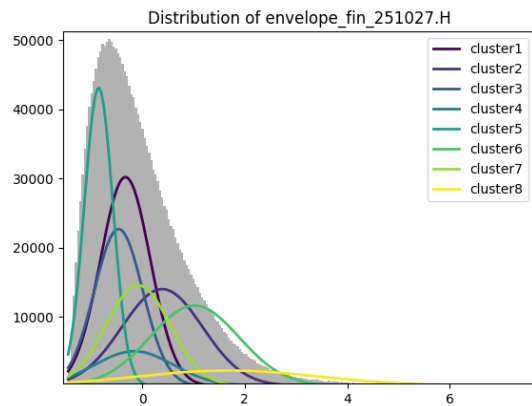
5 References

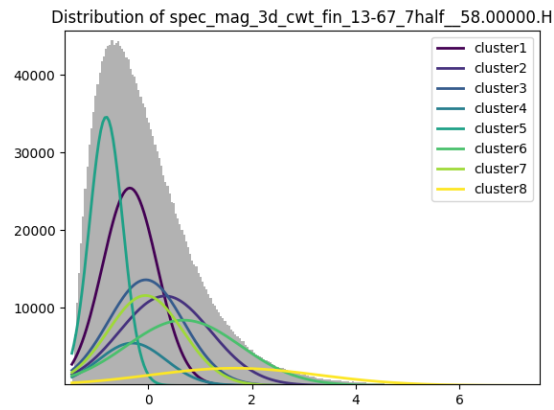
- AASPI Consortium. (2023, October 16). *Geometric Attributes: Computing Texture Attributes – Program glcm3d*. Retrieved from Attribute Assisted Seismic Processing and Interpretation: <https://mcee.ou.edu/aaspi/documentation/>
- Chopra, S., & Alexeev, V. (2006). Applications of texture attribute analysis to 3D seismic data. *The Leading Edge*, 934.
- Chopra, S., & Marfurt, K. (2016). Spectral decomposition and spectral balancing of seismic data. *The Leading Edge*, 176.
- Chopra, S., & Marfurt, K. (2024). *Essentials of Seismic Attributes and Impedance Inversion*. Houston: Society of Exploration Geophysicists.
- Coléou, T., Poupon, M., & Azbel, K. (2003, October). Unsupervised seismic facies classification: A review and comparison of techniques and implementation. *The Leading Edge*, 942.
- DownUnder GeoSolutions Pty Ltd. (2013). *Kokako 3D MSS Processing Report*. New Zealand Petroleum & Minerals.
- Franzel, M., & Back, S. (2018). Three-dimensional seismic sedimentology and stratigraphic architecture of prograding clinoforms, central Taranaki Basin, New Zealand. *International Journal of Earth Sciences*, 475.
- Franzel, M., & Back, S. (2019). Correction to: Three-dimensional seismic sedimentology and stratigraphic architecture of prograding clinoforms, central Taranaki Basin, New Zealand. *International Journal of Earth Sciences*, 497.
- Haralick, R., Shanmugam, K., & Dinstein, I. (1973, November). Textural Features for Image Classification. *IEEE Transactions on Systems, Man, and Cybernetics*, 610.
- Kohonen, T. (1982). Self-organized formation of topologically correct feature maps. *Biological Cybernetics*, 59.
- La Marca, K., & Bedle, H. (2022, November). Deepwater seismic facies and architectural element interpretation aided with unsupervised machine learning techniques: Taranaki basin, New Zealand. *Marine and Petroleum Geology*(136).
- La Marca, K., Bedle, H., Stright, L., & Marfurt, K. (2024, August). Uncertainty assessment in unsupervised machine-learning methods for deepwater channel seismic facies using outcrop-derived 3D models and synthetic seismic data. *Interpretation*, T303.
- Lubo-Robles, D., Devegowda, D., Jayaram, V., Bedle, H., Marfurt, K., & Pranter, M. (2020). Machine Learning model interpretability using SHAP values: application to a seismic facies classification task. *SEG International Exposition and 90th Annual Meeting* (p. 1460). Houston: The Society of Exploration Geophysicists.

- Lundberg, S., & Lee, S.-I. (2017). A Unified Approach to Interpreting Model. *Neural Information Processing Systems* (p. 4765). Long Beach: Curran Associates, Inc.
- Marfurt, K. (2006, July-August). Robust estimates of 3D reflector dip and azimuth. *GEOPHYSICS*, 71, 29.
- Marfurt, K., Kirlin, R., Farmer, S., & Bahorich, M. (1998). 3-D seismic attributes using a semblance-based coherency algorithm. *Geophysics*, 1150.
- McLachlan, G., & Peel, D. (2000). *Finite Mixture Models*. New York: Wiley.
- Roberts, A. (2001). Curvature attributes and their application to 3D interpreted horizons. *First Break*, 85.
- Weimer, P., & Slatt, R. M. (2007). *Introduction to Petroleum Geology of Deepwater Settings*. Tulsa: AAPG.
- Zhao, T., Zhang, J., Li, F., & Marfurt, K. (2016, February). Characterizing a turbidite system in Canterbury Basin, New Zealand, using seismic attributes and distance-preserving self-organizing maps. *Interpretation*, SB79.

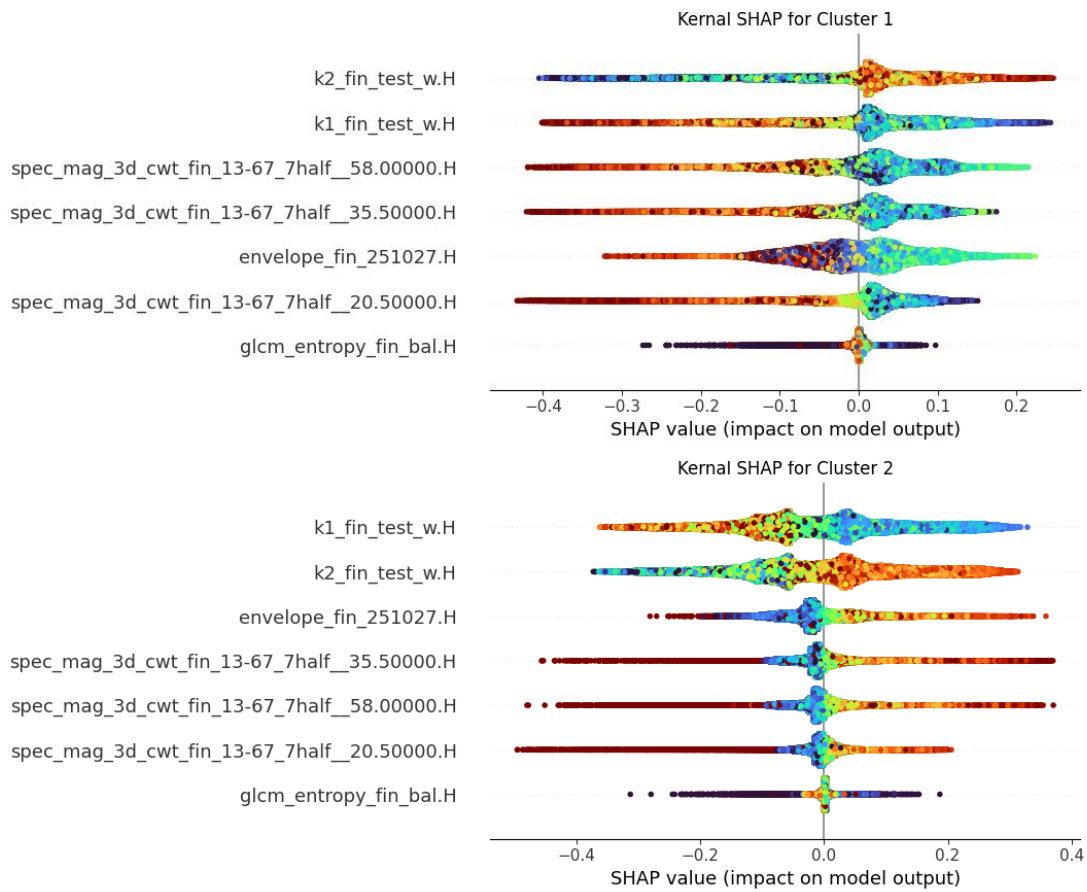
Appendix

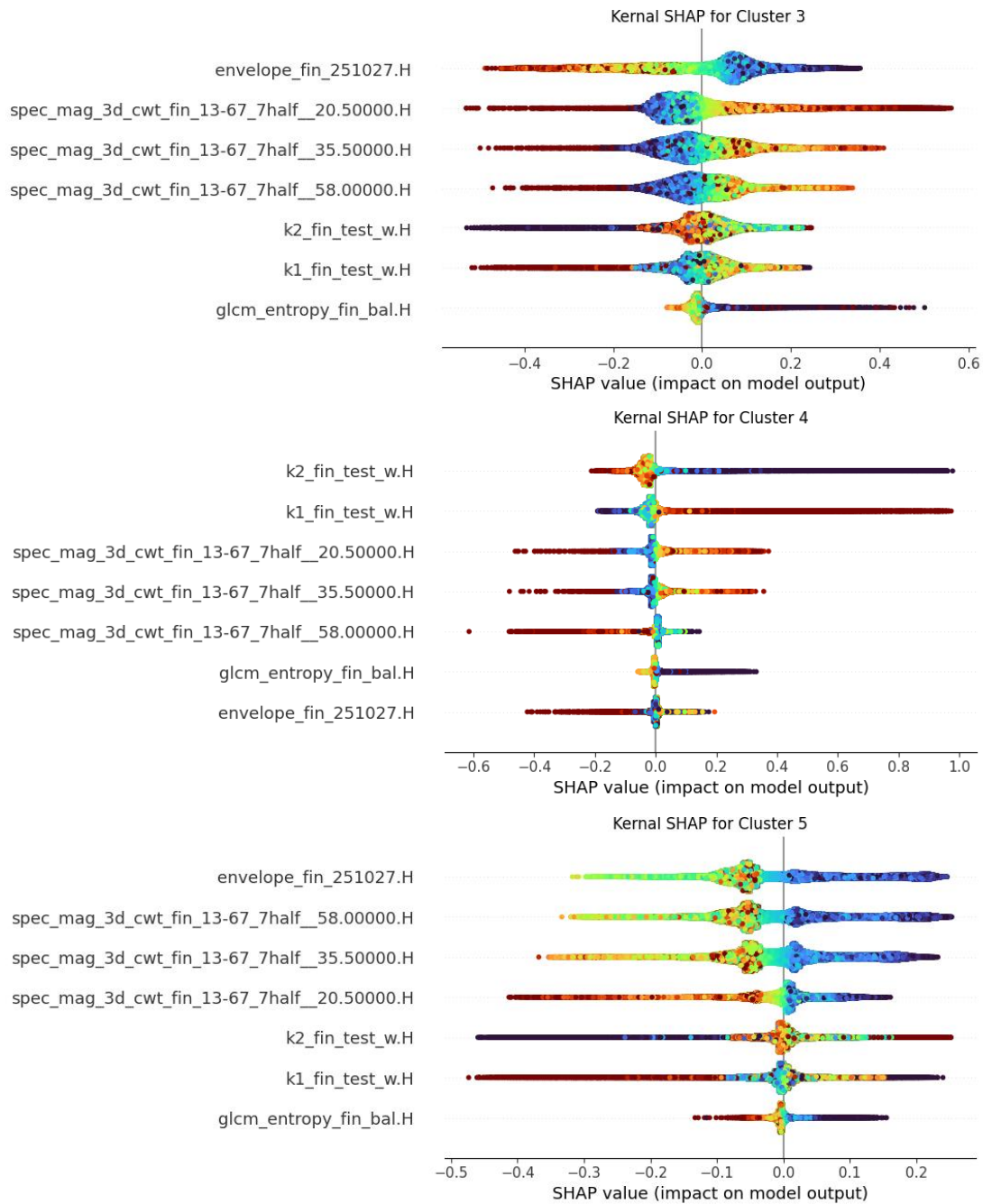
A.1. Histogram and Cluster Gaussians

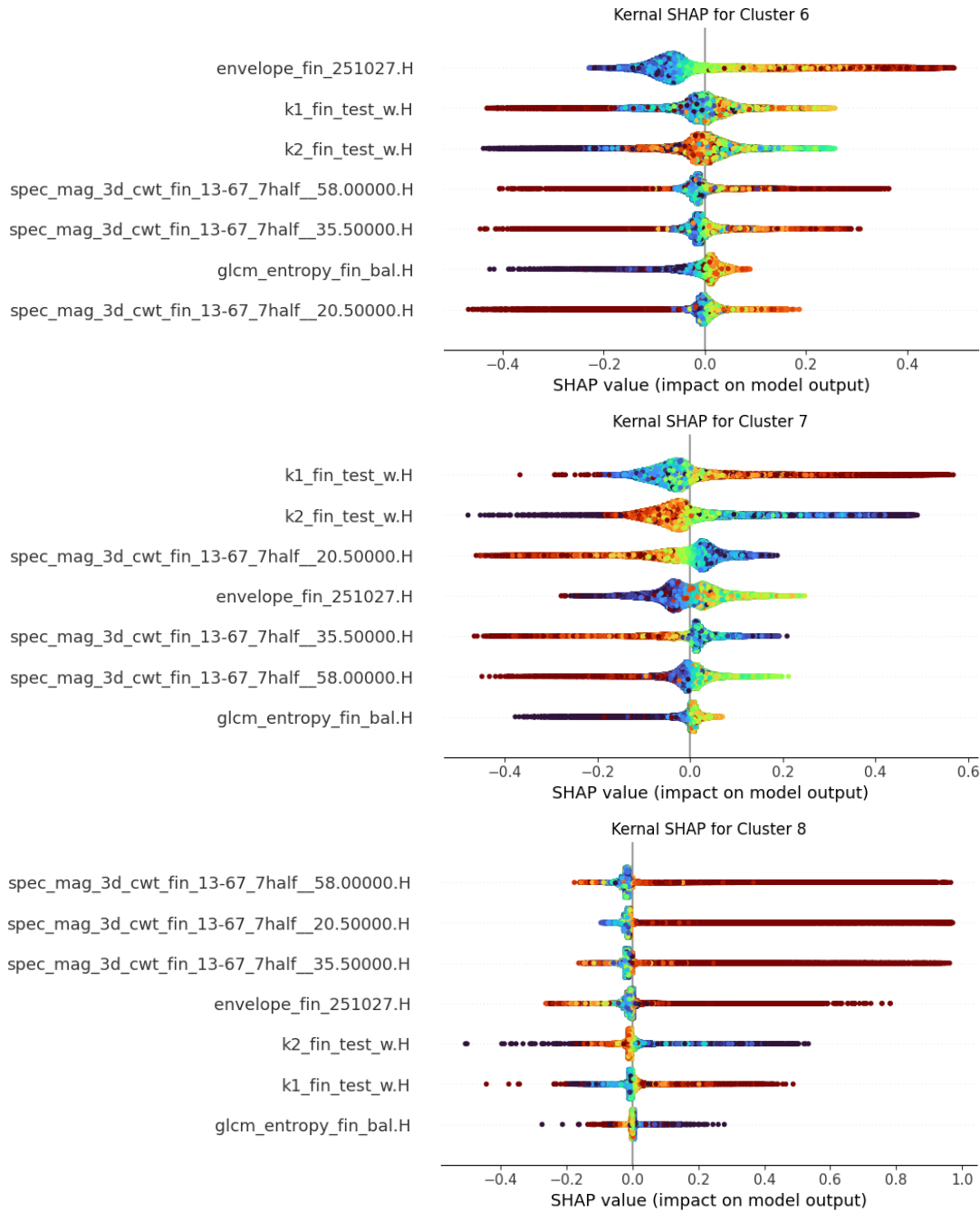




A.2. SHAP graphs



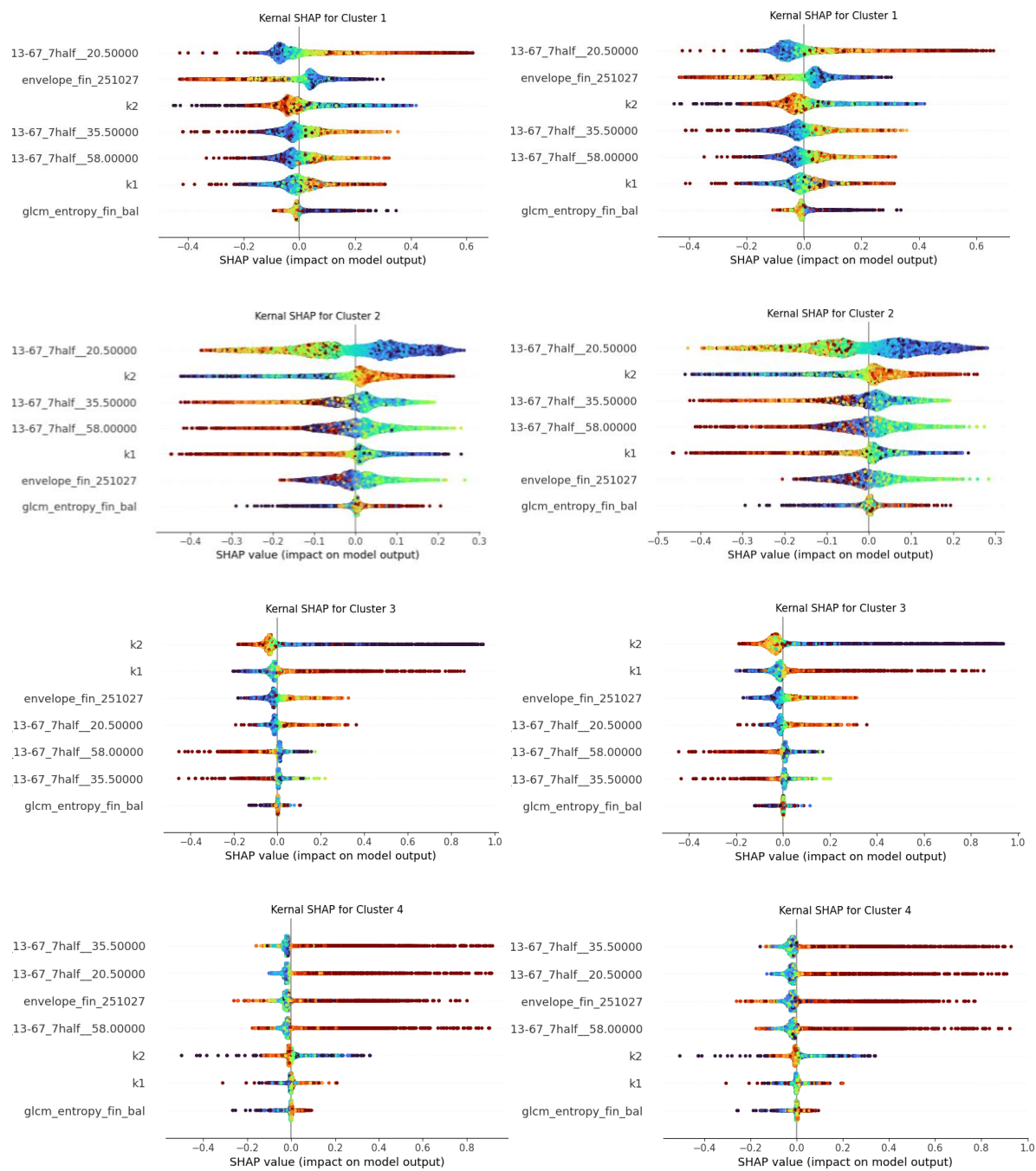


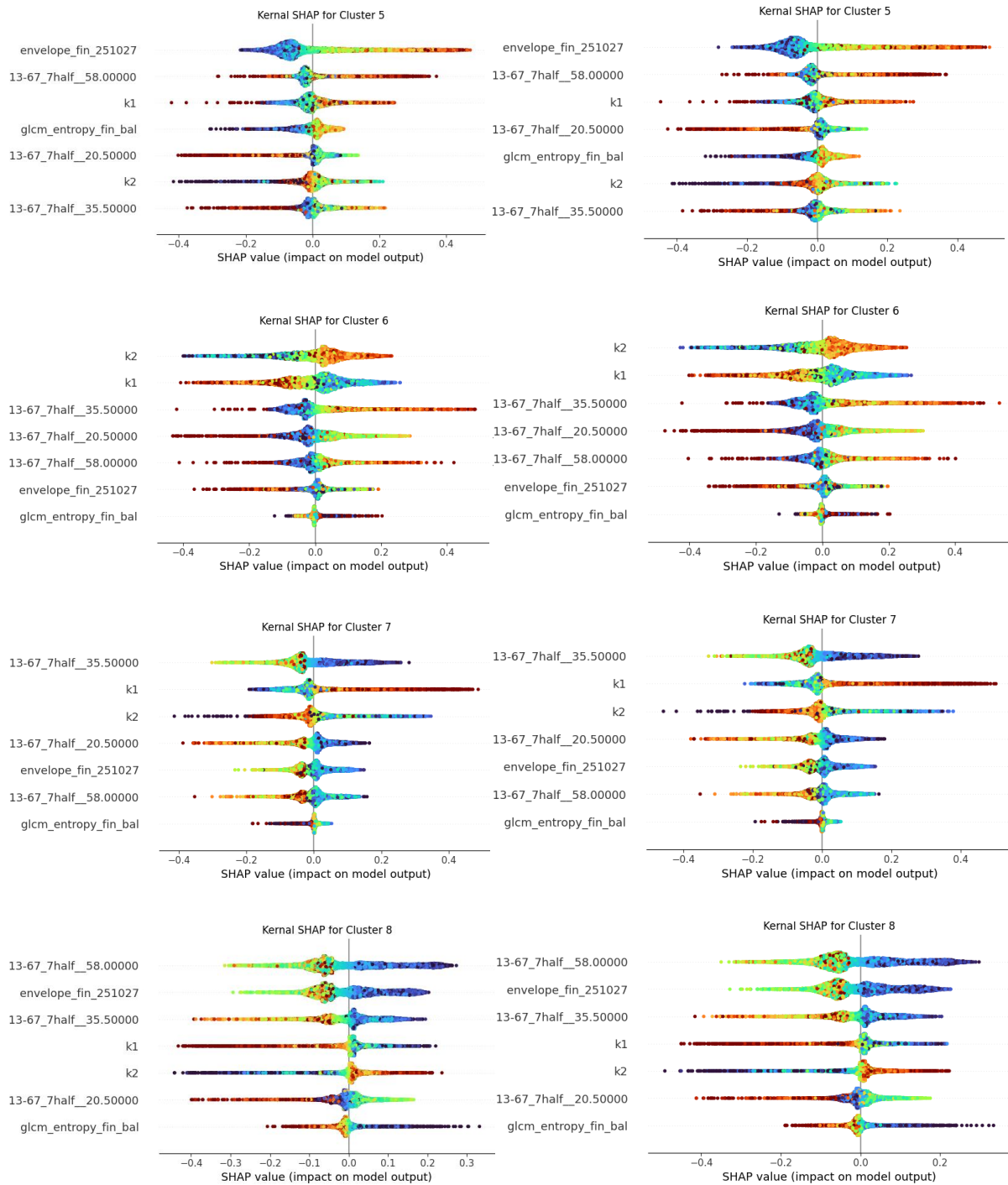


A.3. Batched Kernel SHAP

Due to its cross-sampling nature of Kernel SHAP in approximating the true Shapley value (for each datapoint it goes through the background dataset to simulate absence), processing time grows exponentially as the number of observed data points grows. To combat this inefficiency, we tested a random sampling batch approach in hopes of converging to the same SHAP conclusions but with a significantly faster time by leveraging multicore processing.

We process batches of data points and accumulate them into a single figure covering the whole dataset. When compared to computing the whole dataset at once we end up with practically the same SHAP results:

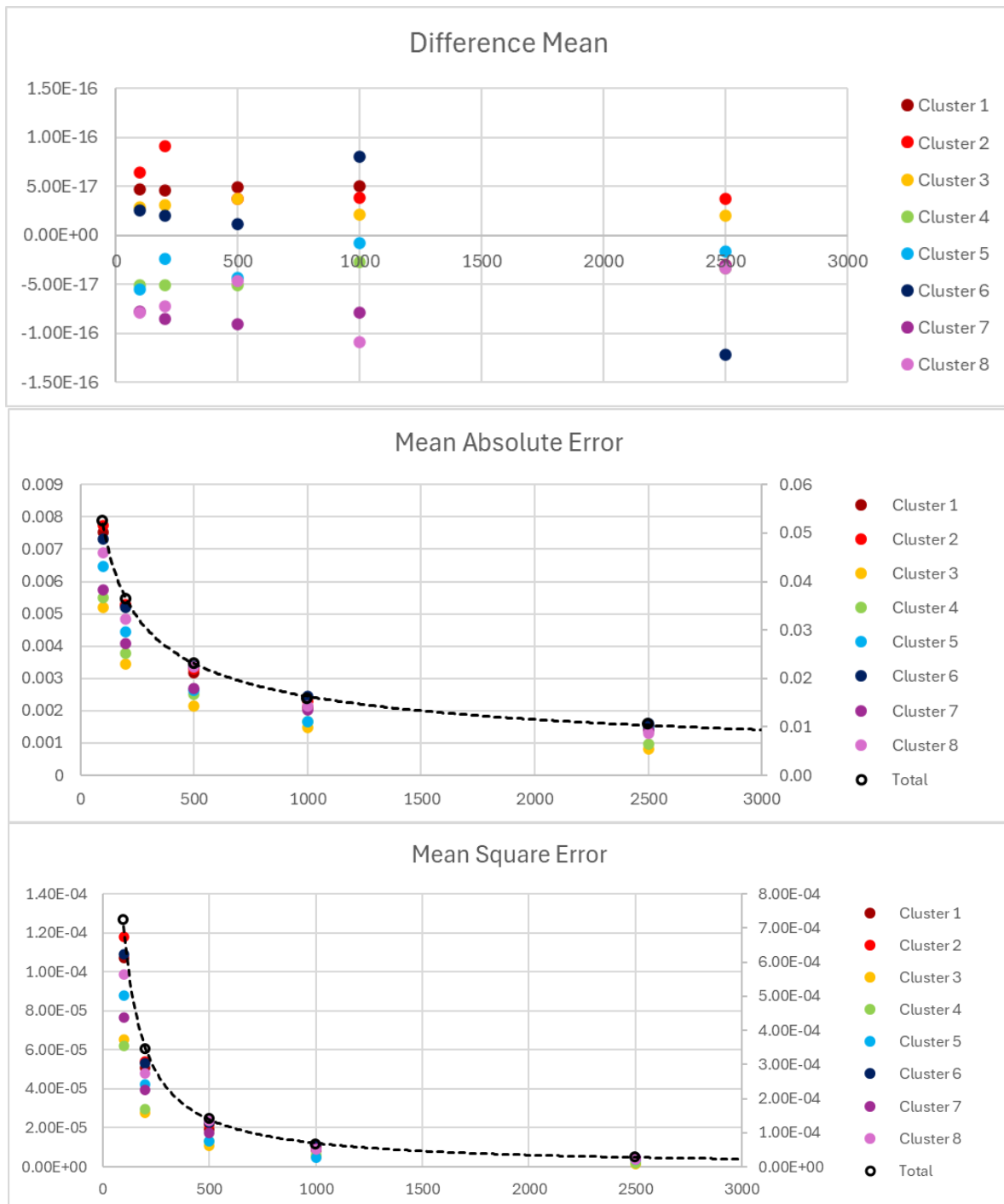




full 10000 data points kernel SHAP results (left) vs 100 datapoint batches (right). Note that the summary result is visually the same, ultimately coming to the same graphical description.

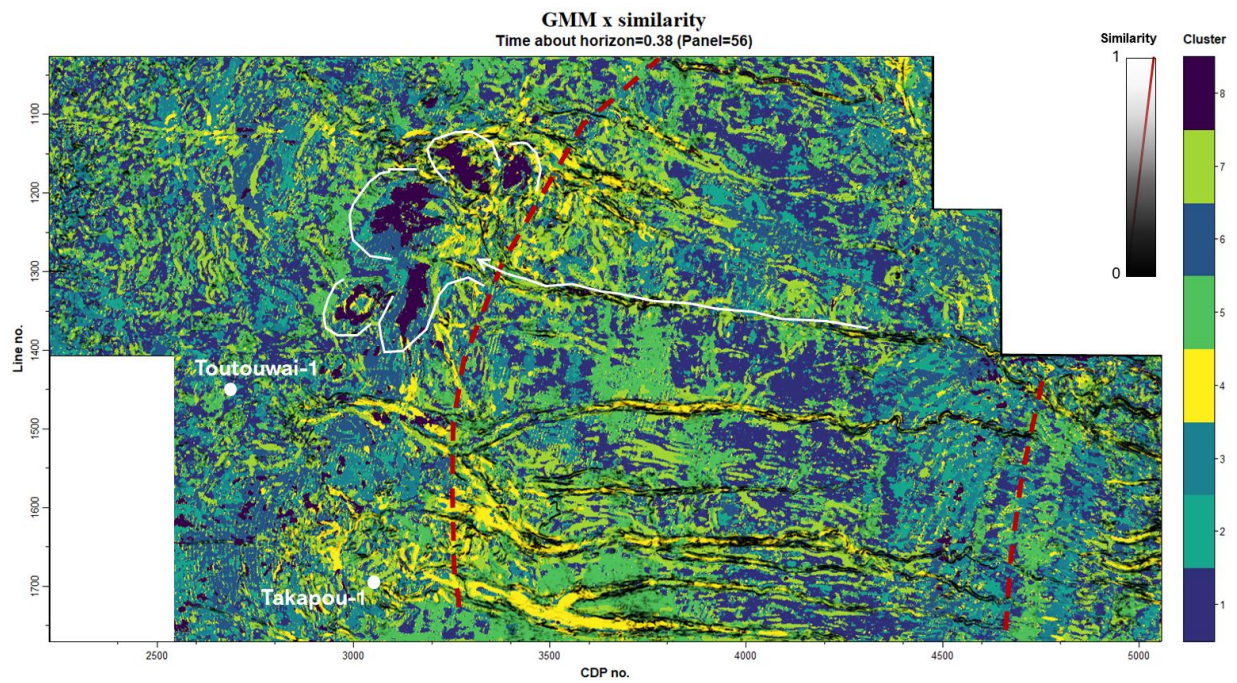
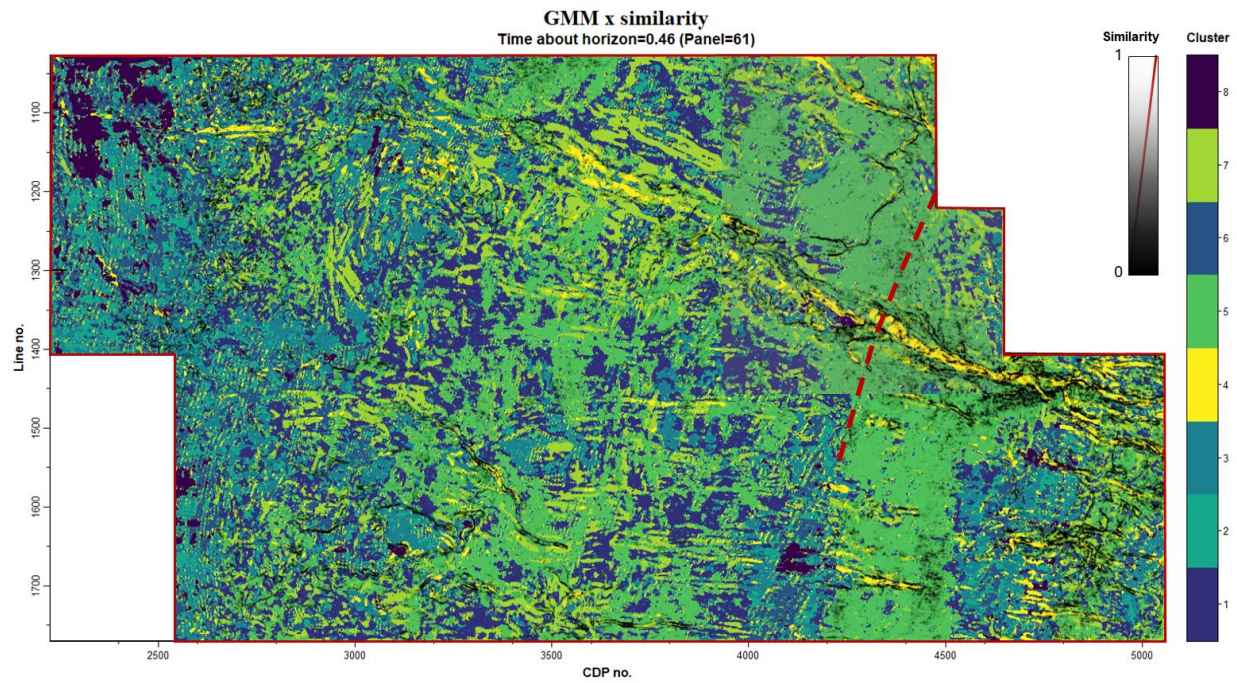
The batch processing (100 data points) took significantly faster with each batch only taking 4 seconds, meanwhile the full dataset (10000) took 10 hours. For additional reference, 200 data points batch took 17 seconds each batch, and 500 data points took 1 minute 48 seconds. This is computed

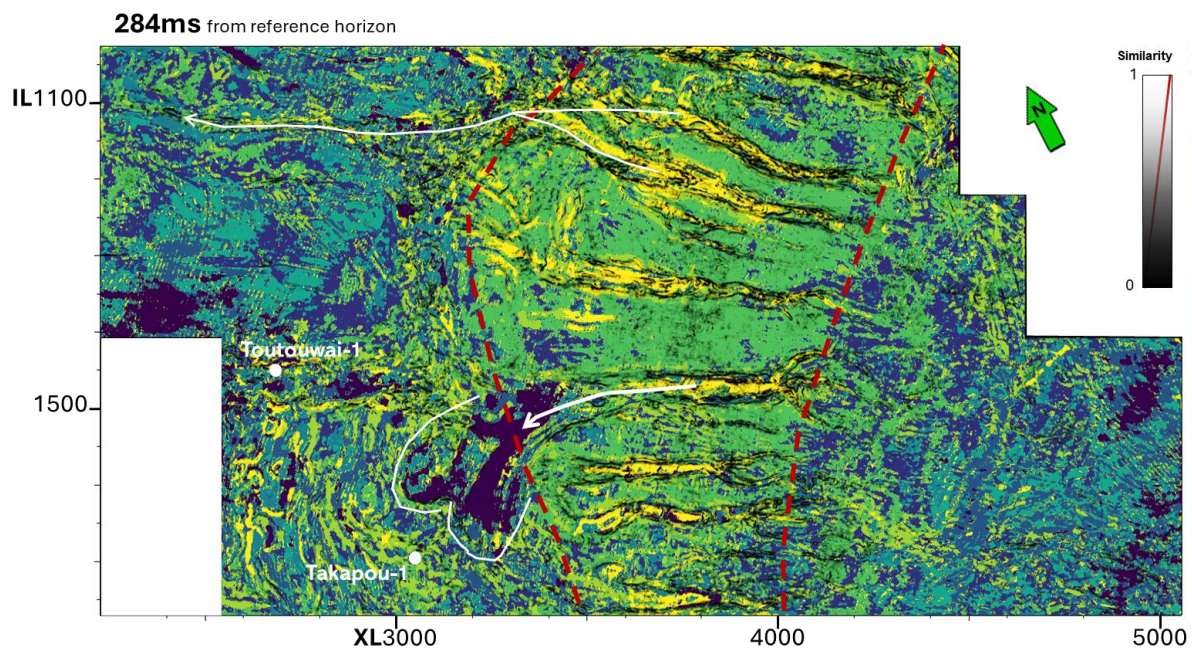
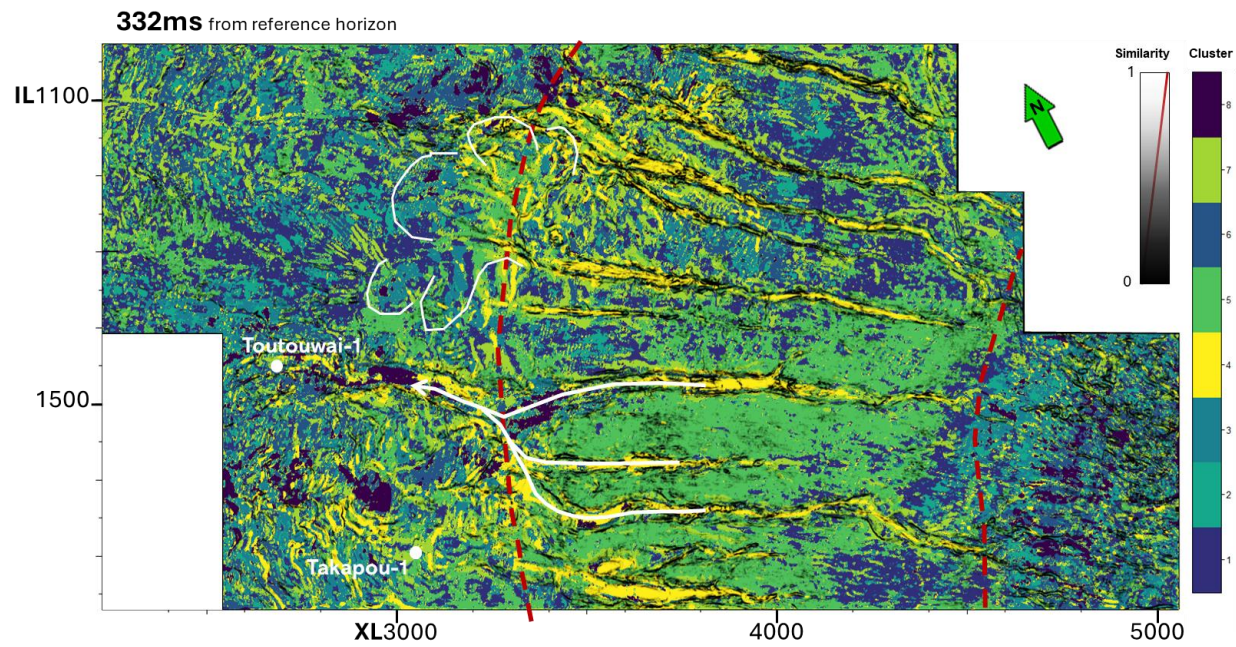
using Windows 11 machines with AMD Ryzen Threadripper PRO 5965WX: 3.8GHz clock speed, and 128GB memory.

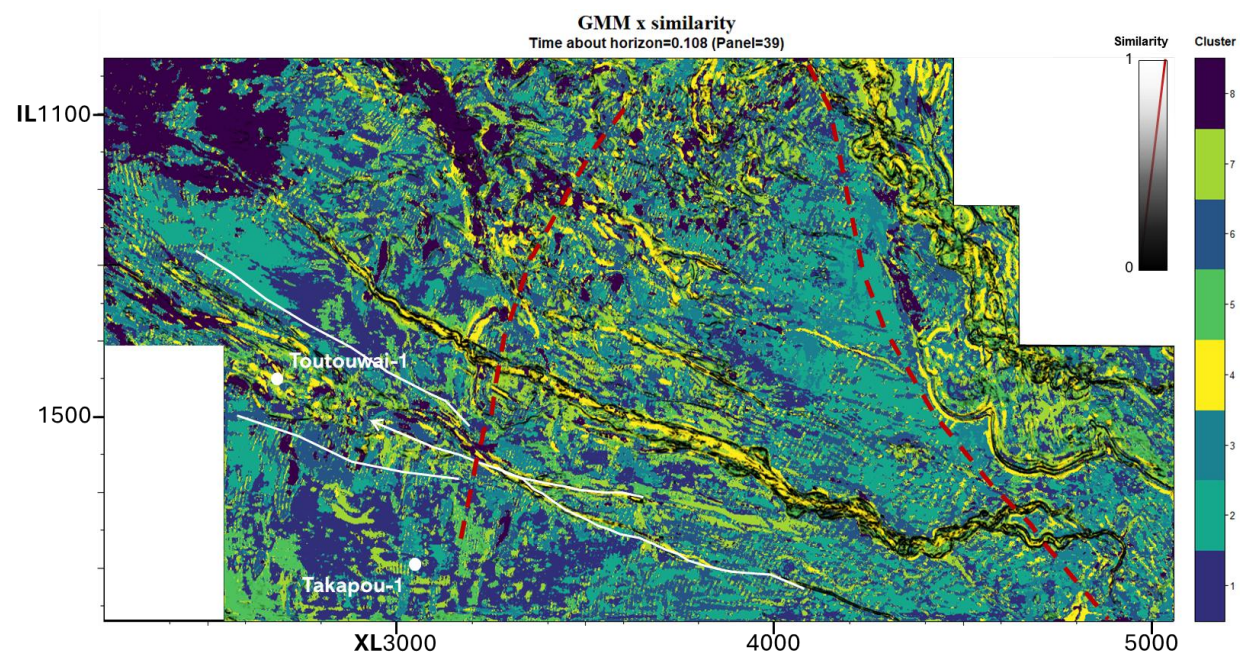
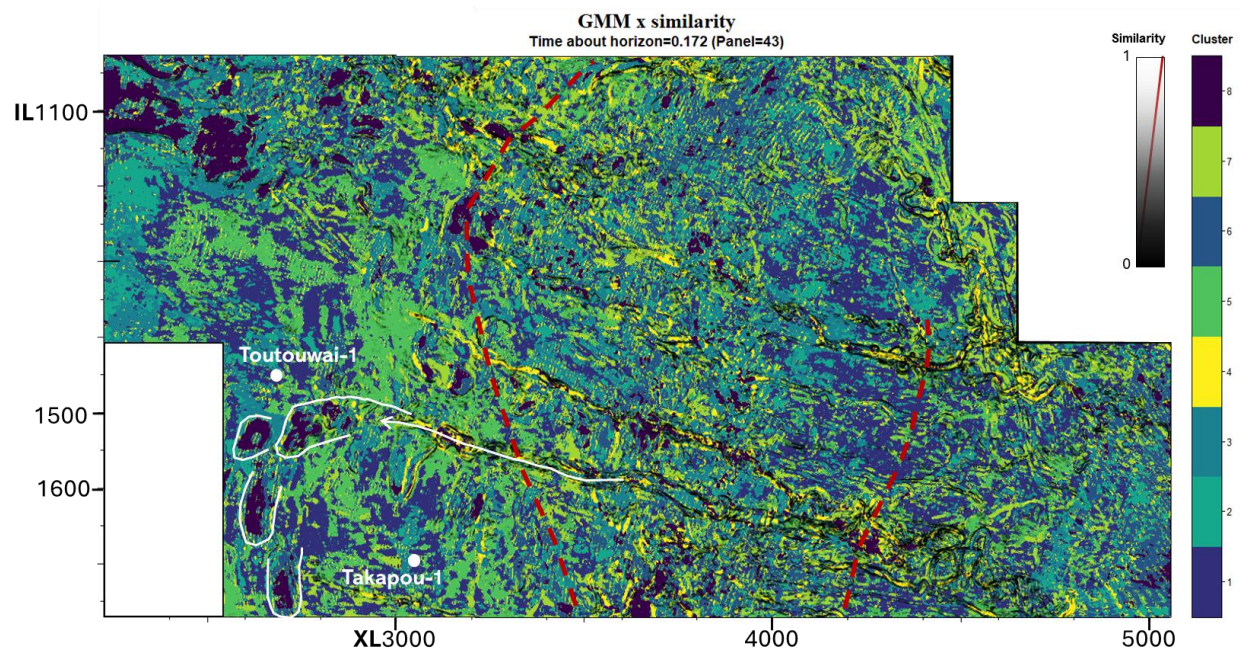


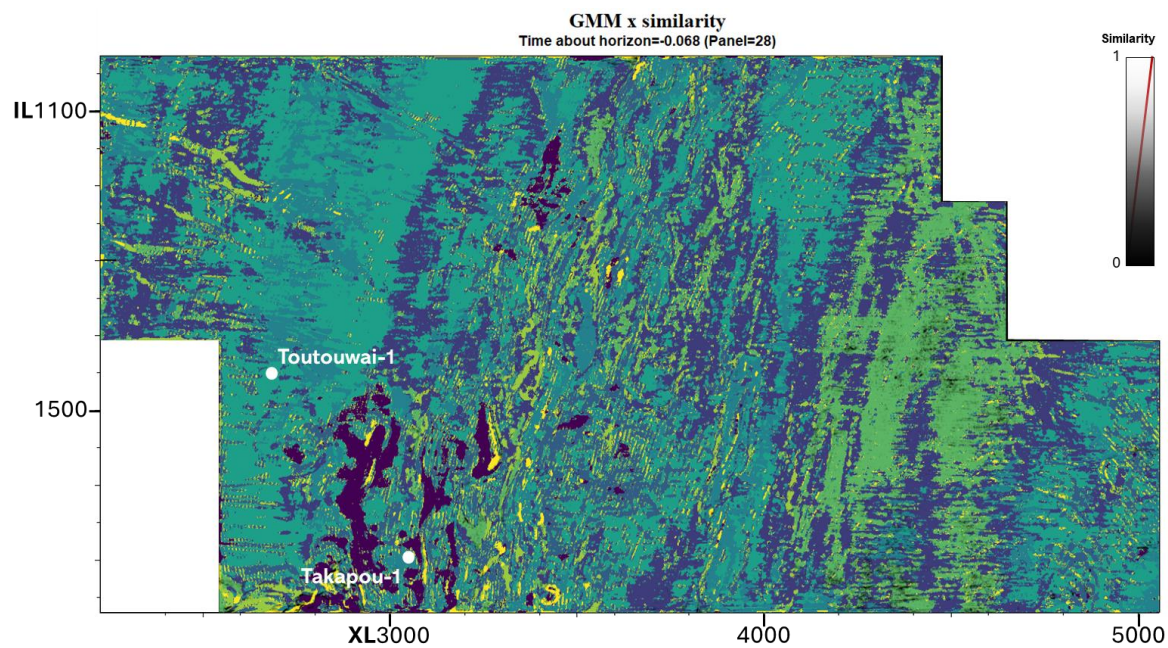
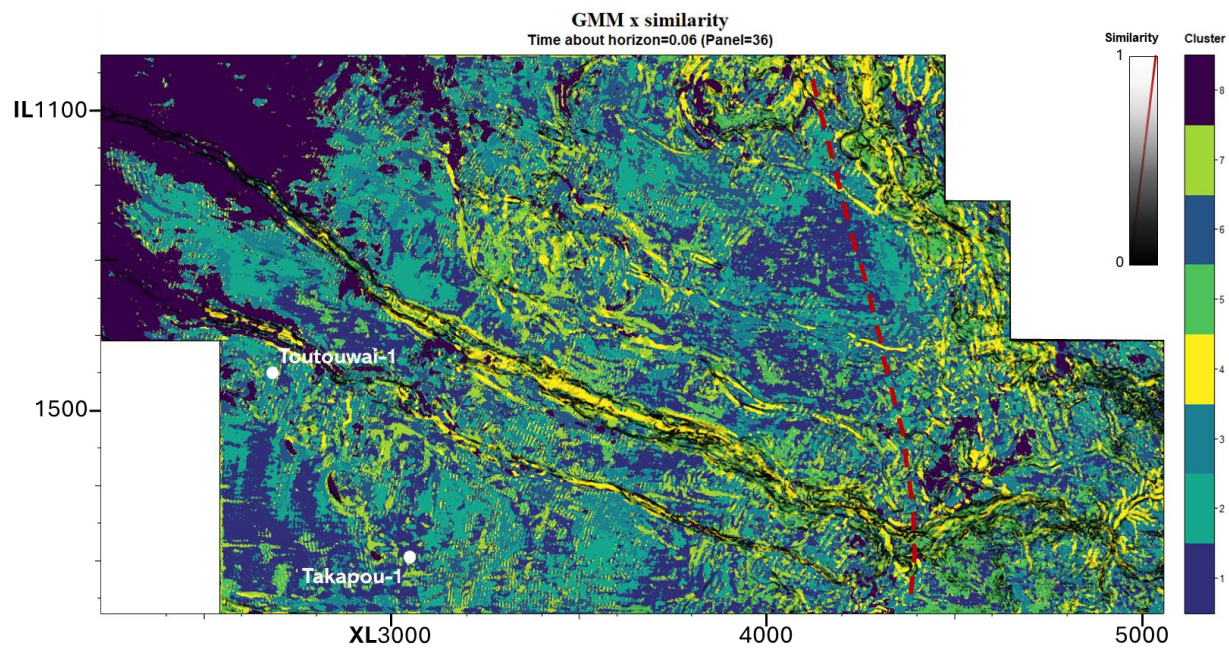
Error calculation comparing full 10000 data points Shapley values and using Batched Kernel SHAP of varying batch sizes (x-axis).

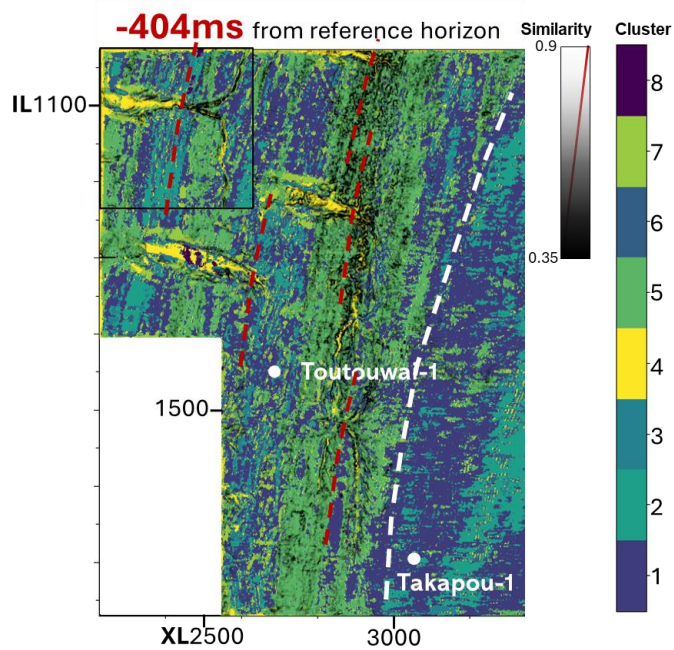
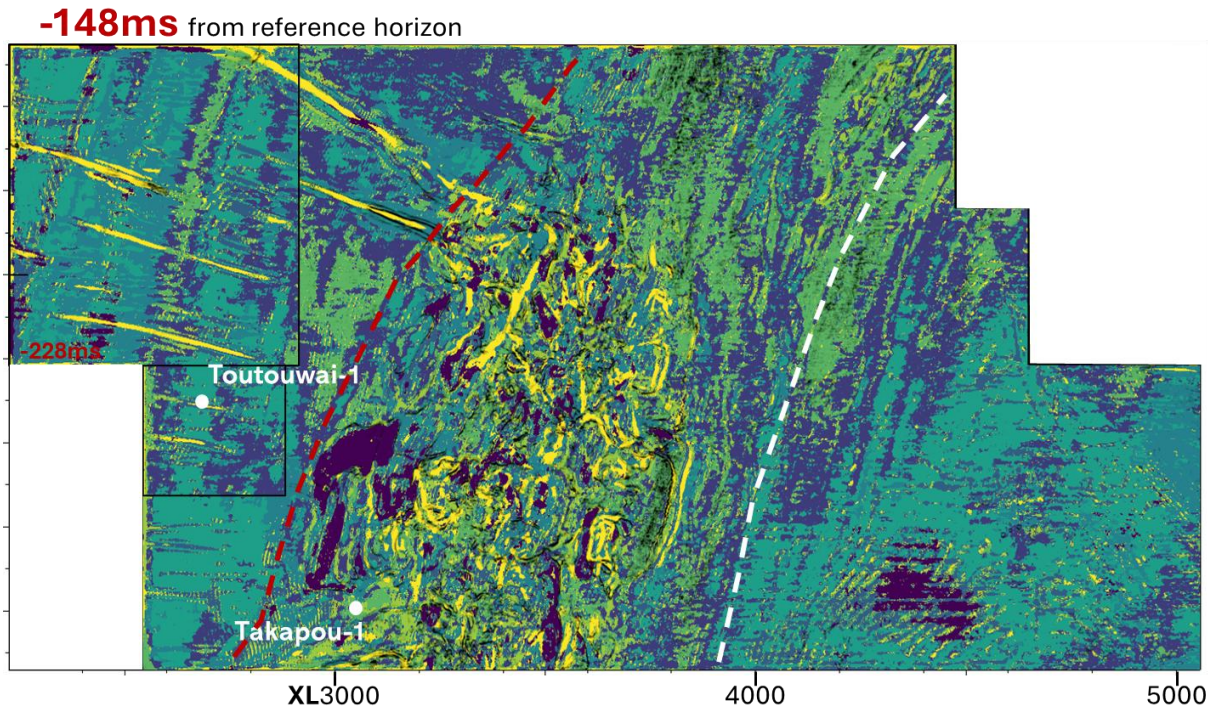
A.4. GMM Maps (from oldest-youngest horizon)







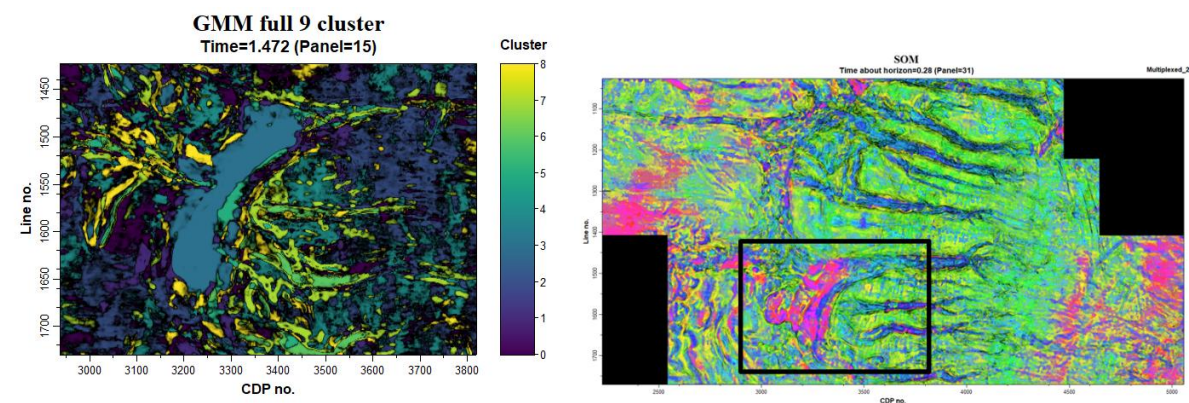




A.5. Preliminary

Our preliminary results are from applying the general workflow to the cropped ‘prototype’ dataset used for algorithm testing with significantly less area of coverage, a limited time window and

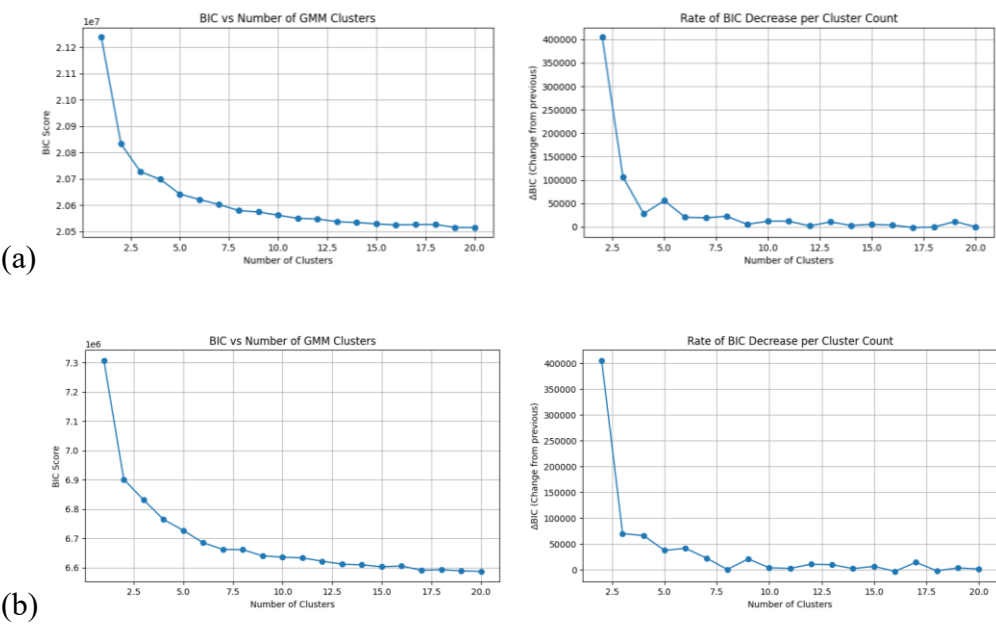
expected less geological complexity: mainly a submarine fan structure and its slope channels as seen on the quick SOM (Self Organizing Map) machine learning volumetric classification.



A.6. Training dataset dependence

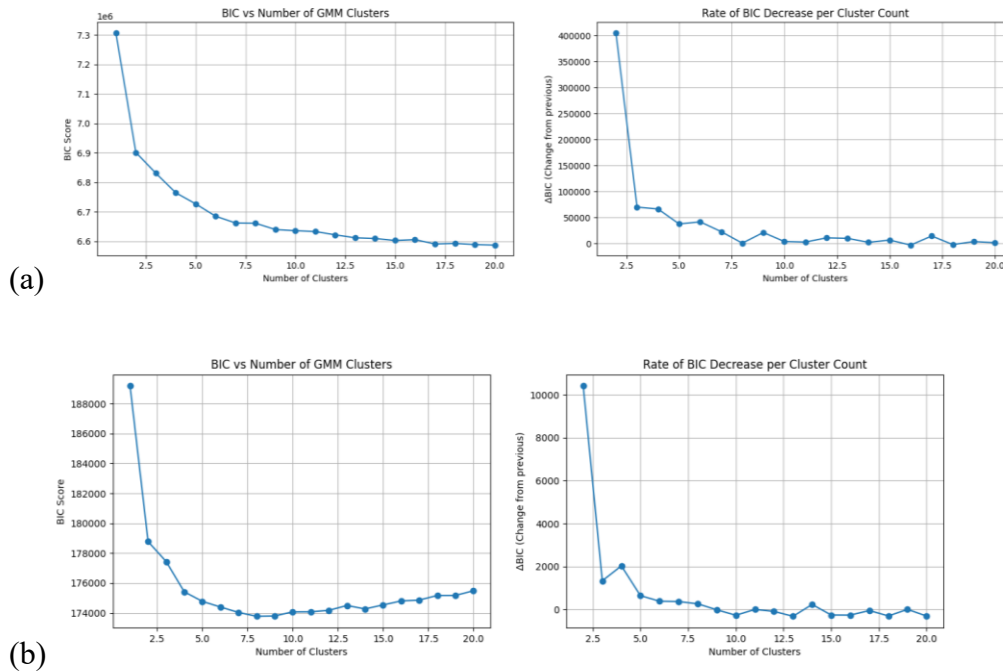
As an unsupervised machine learning method, GMM is sensitive to its input dataset. Careful normalization and regularization (controlled sampling) of the dataset is crucial to have a consistent result in capturing the desired signal variation, while keeping an efficient processing timeline.

We tested the impact of common normalization practices (z-score, robust, logarithmic) for multi-attribute machine learning and how GMM performs in finding the fit with its probability density functions, compared to without normalization (Gaussian normal distributions):



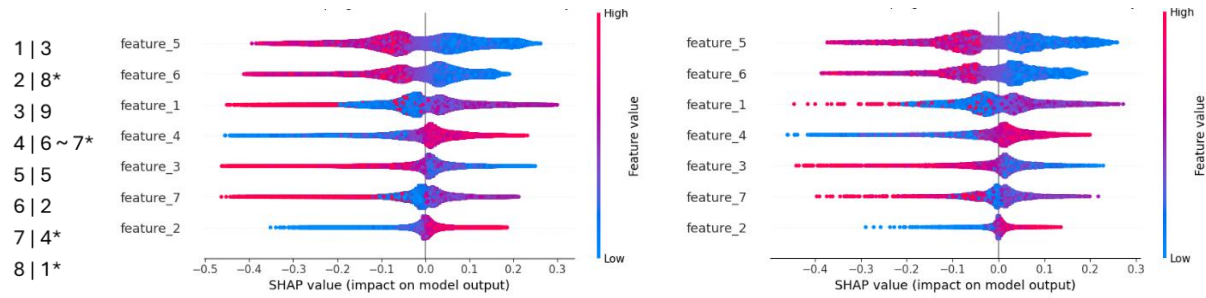
BIC analysis of (a) non-scaled and (b) Z-score scaled dataset

Even with the *same* data points, the Z-score scaled dataset has a completely different evolution and significantly lower BIC score, implying an inherent change in how the algorithm finds convergence. This confirms that the GMM algorithm, although mainly pays attention to the distribution, is not independent of the data values, and thus proper data pre-conditioning will be beneficial to its performance. To further improve computing efficiency, we also tested the impact of regularizing/decimating the dataset by resampling it to 16IL x 32XL x 20ms sampling:



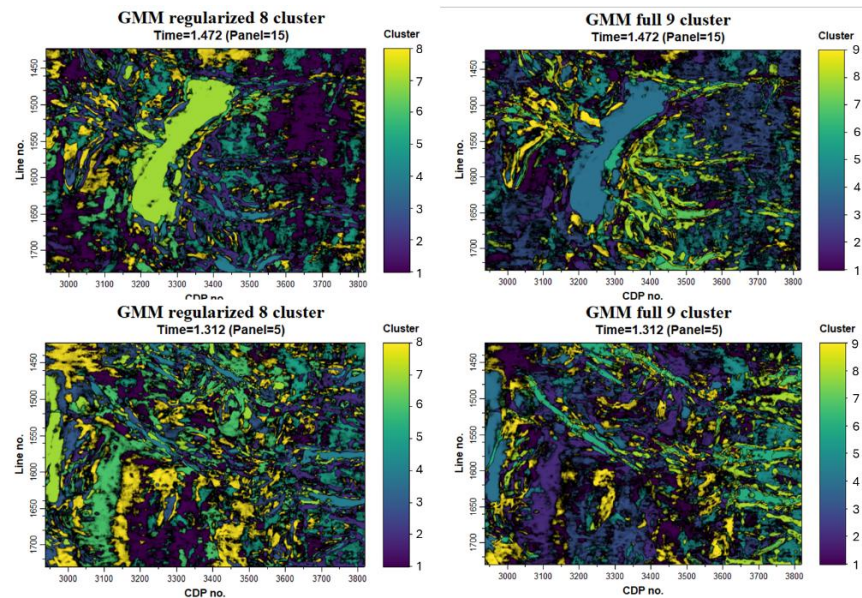
BIC analysis of non-regularized (a) and regularized (b) training dataset.

We find that BIC value converges faster on the regularized dataset and exhibits an increase past a certain cluster number (8 clusters) indicating where the model starts to become redundant/overfitting. The non-regularized dataset does not have a clear indication of where overfitting begins and shows a constant decrease of BIC value past a certain point, suggesting that the added Gaussian clusters are able to fit certain background distributions and edge cases. This makes sense as the non-regularized dataset could pick up finer details but also picks up more background noise as consequence.



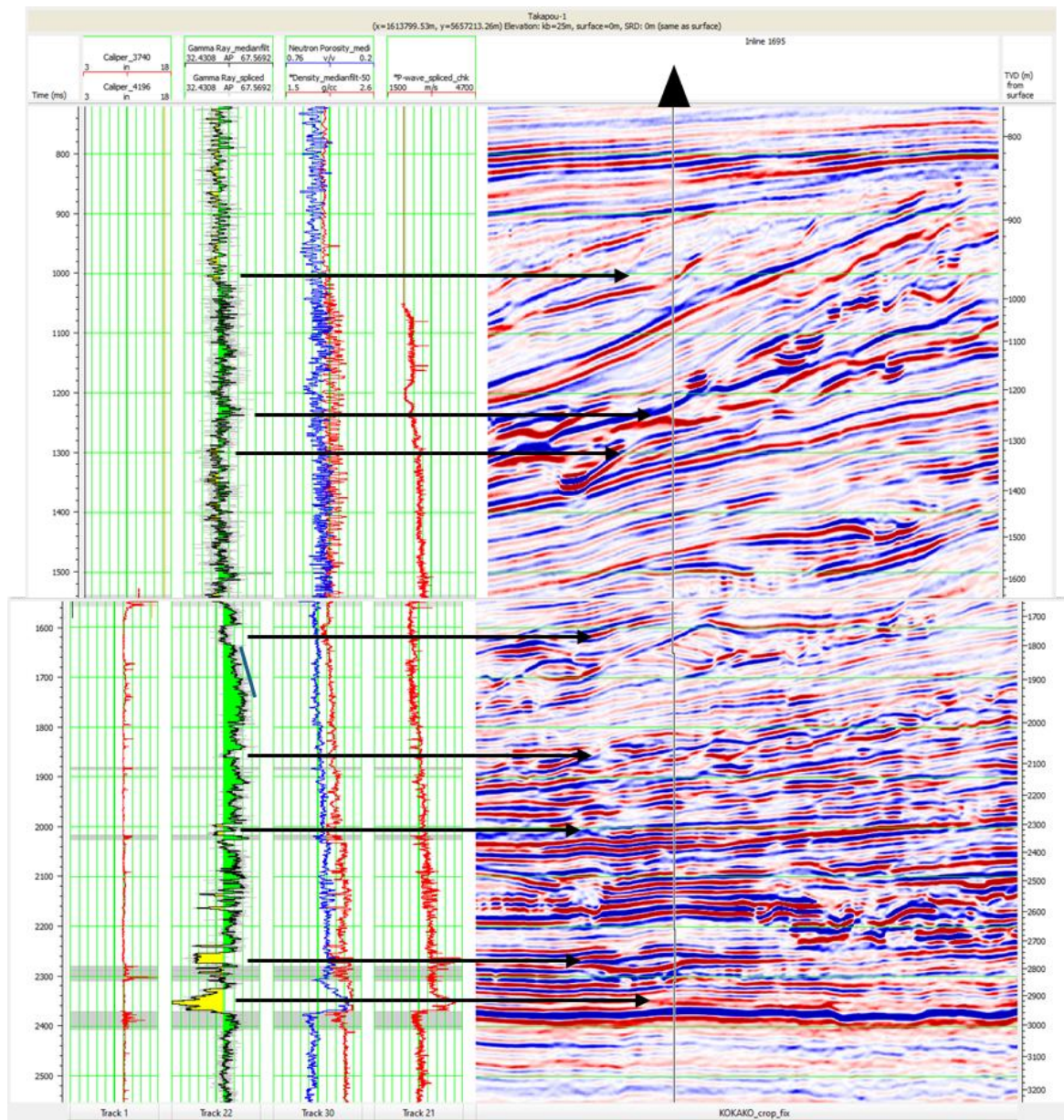
SHAP beeswarm plots of the same cluster facies from non-regularized dataset (left), and regularized dataset (right). This specific cluster was able to be identified in the same way indicating that this signal was sampled by both datasets.

The higher number of data points would mean longer processing time, especially during intensive Kernel SHAP processing. Picking a balance between finer sampling and processing time is important as we note that the non-regularized dataset, which amounts to 389000 data points, took 30 hours to complete while the regularized dataset with only 10000 data points took 45 minutes (performed on a single core for consistency); and this is on a cropped prototyping dataset.

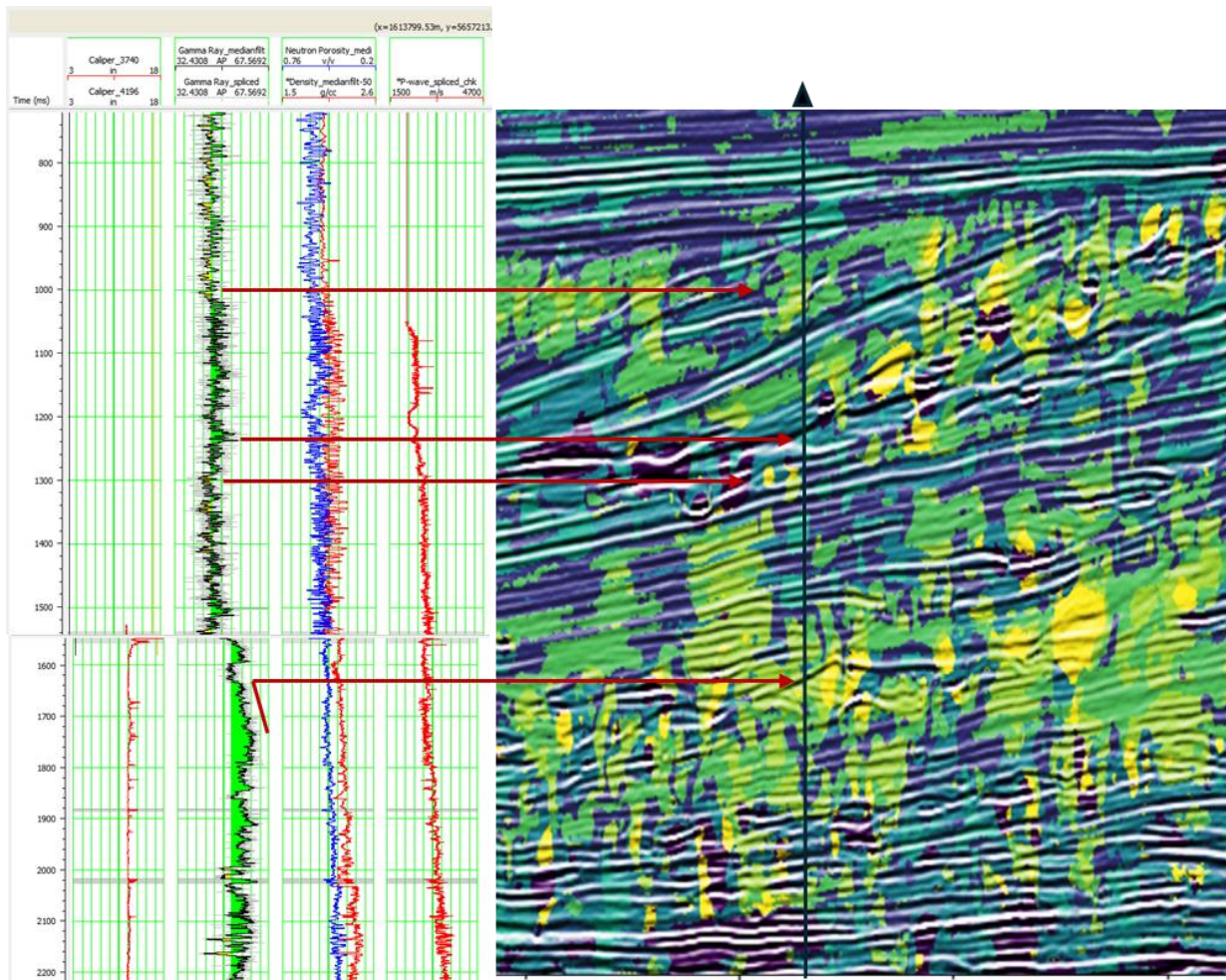


regularized 8 cluster GMM result (left), versus full dataset 9 cluster GMM (right)

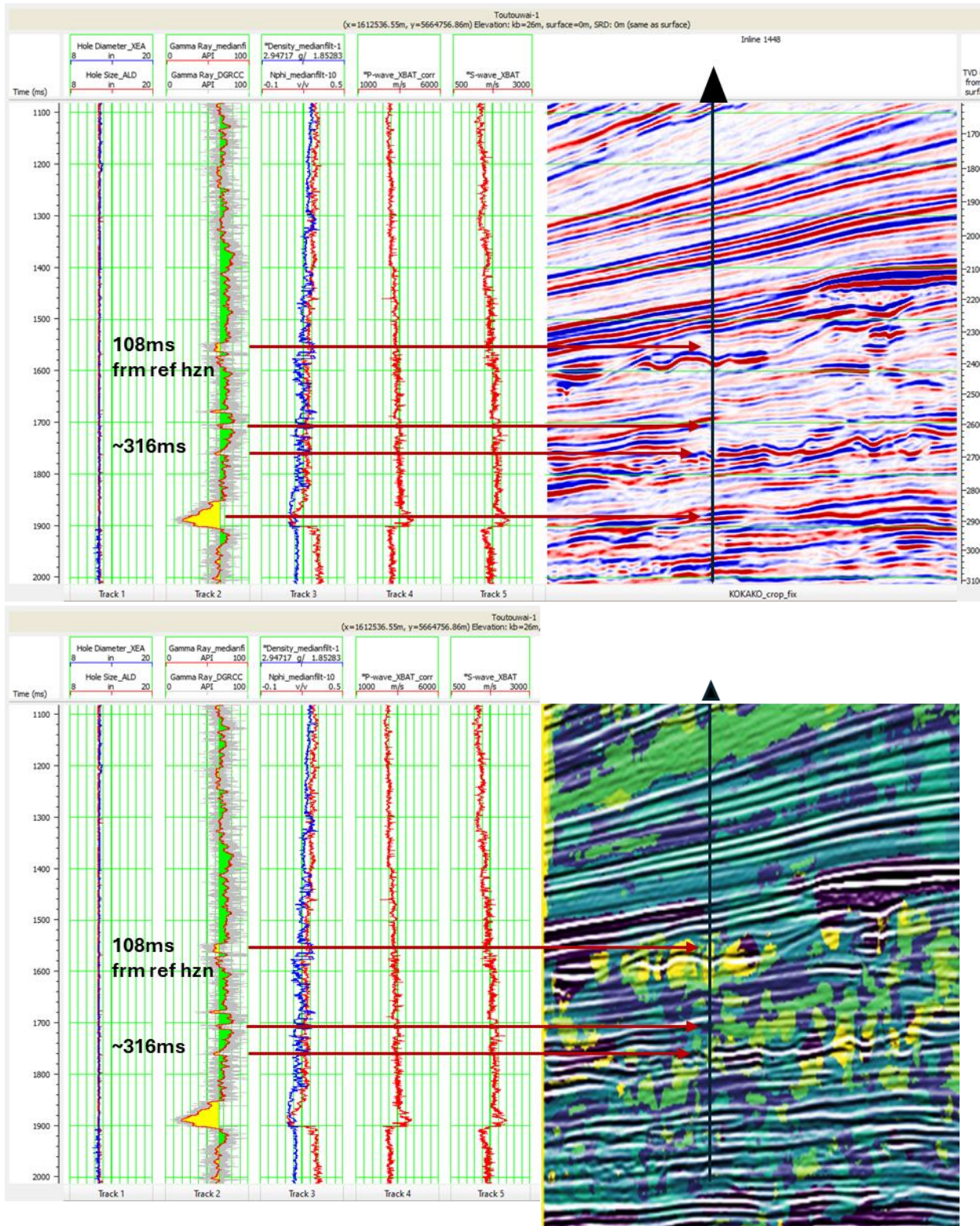
A.7. Well Logs



a) Takapou-1 log and well path on the seismic.



b)



Toutouai-1 is positioned better than Takapou-1 as it successfully hit the target facies. Unfortunately the log data is only LWD (Logging While Drilling) log which has a lower quality (noise) and especially the sonic log. GR log is good.