

DETERMINING THE IMPACT OF THE HUMAN OPERATOR FACTOR USING A NAVIGATION SIMULATOR

Kochubei P., Graduate student, Kherson State Maritime Academy, Ukraine, e-mail: pavlo.kochubei@ksma.ks.ua;

Nosov P., Ph.D., Associate Professor, Department of ship computer systems and networks, Kherson State Maritime Academy, Ukraine, e-mail: pason@ukr.net, ORCID: 0000-0002-5067-9766.

As maritime operations become increasingly complex, there is a growing need for objective, scalable methods to evaluate seafarer competence beyond traditional instructor-based assessments. This study presents a comprehensive, data-driven framework for analyzing cadet performance in maritime simulator exercises, utilizing state-of-the-art unsupervised learning and explainable AI. The research draws on multivariate time-series data from navigational simulations, capturing vessel dynamics, control actions, and environmental parameters across dozens of features. A rigorous preprocessing pipeline was applied, combining statistical feature aggregation and redundancy reduction through Pearson correlation and Mutual Information. This yielded a compact, yet informative feature set encompassing control inputs, navigational states, and vessel motions. To offset limited number of sessions small in the dataset, each simulation was split into meaningful intervals by applying rolling window statistics. Each window was then encoded as a summary vector, reflecting both central tendencies and temporal variability. The analysis employed the HDBSCAN clustering algorithm, which excels in detecting groups of variable density and naturally identifies outlier behaviors—critical in the context of training evaluation. The resulting clusters were projected into lower-dimensional space via t-SNE, providing interpretable visualizations of cadet performance patterns. To further elicit the distinguishing characteristics of each group, a linear Support Vector Machine was trained to predict cluster membership, with SHapley Additive exPlanations (SHAP) attributing each decision to underlying features. Key findings reveal that clusters align with distinct navigational strategies: stable, conservative approaches are differentiated from more dynamic or risk-prone styles by features such as roll velocity, yaw rate, and engine RPM. Sessions flagged as outliers typically exhibited abrupt maneuvers or inconsistent control usage, highlighting potential skill gaps. The SHAP-based interpretability layer transforms complex model outputs into actionable instructional feedback, enabling targeted interventions and tailored training. Overall, this automated approach has potential to become a transparent, scalable alternative to subjective grading in maritime education, with significant implications for enhancing safety and developing individualized learning pathways. The proposed system demonstrates strong potential for integration into real-world training environments and continuous improvement as more operational data becomes available.

Key words: water transport; operation of transport facilities; navigation safety; human factor; automation; risk; intelligent systems; LAD.

DOI: 10.33815/2313-4763.2025.1.30.158-170

Introduction. Despite technological advances in the maritime industry – ship design, navigational systems, and sensor technologies – human error remains a primary contributing factor to naval accidents, accounting for more than 85% of all accidents [1]. The International Maritime Organization (IMO) continues to emphasize the critical role of human factors in maritime safety, especially in complex navigational contexts such as port approaches and congested sea lanes [2].

Advancements in sensor technologies, such as ECG, eye tracking, etc., increased connectivity of ships, a surge in computational and data collection capabilities, and popularization of machine learning techniques have expanded existing and allowed new capabilities in monitoring and assisting humans in the maritime context [3]. These technological improvements allowed advancements in navigator fatigue detection, advanced trajectory prediction, anomaly detection in AIS data, and many other use cases [4–6].

In parallel, a shortage of qualified training professionals in the maritime labor market causes accelerated promotion of marine professionals. Therefore, cadets have less time to develop necessary seafaring skills. According to a 2021 report by the Baltic and International Maritime Council (BIMCO) and the International Chamber of Shipping (ICS), the global merchant fleet is expected to keep expanding, maintaining a high demand for skilled maritime professionals. Nevertheless, the industry continues to face a shortage of qualified personnel. Projections estimate

that an additional 17,902 officers will be needed annually through 2026 to meet the operational needs of the global merchant fleet [7].

Vocational education must ensure individuals develop the essential knowledge, skills, and competencies required for professional success [8]. Therefore, it's necessary to create effective and sound methods of preparing maritime staff. The maritime industry has long been relying on simulator training to optimize education. However, assessment methods in simulation training are predominantly subjective, relying primarily on the judgment and experience of instructors. This introduces variability and potential bias in the evaluation of navigation competencies.

To address these challenges, the application of machine learning to maritime simulator log data offers a promising opportunity. Learning Analytics Dashboards (LADs) have demonstrated value in educational technology domains, offering a framework for integrating objective, data-driven insights into maritime training. Together with predictive models, LADs can improve the feedback for trainees and instructors and serve as an early-warning system for suboptimal performance [9].

Research Purpose and Objectives. The purpose of this study is to develop a method for identifying and evaluating cadet navigational behavior and decision-making strategies in maritime simulator training environments using intelligent, data-driven approaches. The study aims to enhance the objectivity and precision of performance assessment, reduce dependence on instructor subjectivity, and enable early detection of potentially unsafe operational patterns. To achieve this, the following research objectives are formulated:

1. To propose an approach for the automated analysis of navigational and control data recorded during maritime simulation exercises, with a focus on identifying behavioral patterns relevant to operator performance.

2. To establish a framework for classifying and assessing navigational and control strategies, enabling the recognition of outlier behaviors indicative of suboptimal decisions, skill gaps, or elevated operational risk.

3. To examine the potential of unsupervised learning methods for segmenting and evaluating cadet performance to inform adaptive training, targeted feedback, and early-risk intervention strategies in simulator-based education.

Primary Research Material. To achieve the outlined objective, a dataset was prepared that comprises time-series simulation data from 13 independent navigational exercises performed by different cadets using model Navi Trainer Professional 5000 [10] simulator conducting a passage through the Bosphorus strait. The extracted dataset (Fig. 1) captures vessel dynamics and environmental conditions across 44 parameters valid for the exercise under discussion and selected vessel type.

	pos_east	pos_north	yaw	heave	pitch	roll	surge	u_speed	v_speed	rpm_mid	rud_mid	rpm_mid_cmd
0	4.62	13.49	0.000	0.000	0.183	0.000	0.000	0.002	-0.004	0	0	0
1	4.62	13.49	-0.000	-0.001	0.181	0.040	0.001	0.003	-0.004	0	0	0
2	4.62	13.49	0.000	-0.003	0.174	0.032	0.000	0.003	-0.004	0	0	0
3	4.62	13.49	0.000	-0.015	0.166	0.022	0.000	0.003	-0.004	0	0	0
4	4.62	13.49	0.001	-0.011	0.157	0.014	-0.001	0.003	-0.004	0	0	0
5	4.62	13.49	0.001	-0.005	0.150	0.007	-0.000	0.003	-0.004	0	0	0
6	4.62	13.49	0.002	-0.007	0.147	0.002	0.000	0.003	-0.005	0	0	0
7	4.62	13.49	0.003	-0.001	0.147	-0.002	0.001	0.003	-0.005	0	0	0
8	4.62	13.49	0.004	-0.002	0.149	-0.005	0.000	0.003	-0.005	0	0	0
9	4.62	13.49	0.005	0.001	0.154	-0.007	-0.000	0.003	-0.005	0	0	0

Figure 1 – Fragment of data extracted from TRANSAS NTPRO 5000

Here's brief description of the features provided: *pos_east*, *pos_north* are position relative to the starting point of the exercise; *roll* is the tilting motion of a ship from side to side; *pitch* is the up-and-down tilting motion of the ship; *yaw* is side-to-side motion of the ship about its vertical axis; *surge* is the forward and backward motion of the ship along its longitudinal axis; *sway* is the side-to-side motion of the ship along its transverse axis; *heave* is vertical motion of the ship along its vertical axis. *u_speed* forward surge velocity speed, *v_speed* sideways sway velocity, *rpm_mid* indicates Revolutions Per Minute (RPM) of a main propulsion unit; *rud_mid* shows rudder angel and *rpm_mid_cmd* commanded but potentially not yet achieved main engine RPM value.

Upon conducting Exploratory Data Analysis (EDA) [11], the list of meaningful features was reduced to 33 by eliminating non-relevant columns for this vessel type, columns with few values, e.g., the autopilot was turned off for all exercises and therefore had only one value.

Additionally, a set of similarity metrics was computed to avoid supplying duplicate information to the model and further decrease redundant computations – the Pearson Correlation Coefficient [12], which measures the linear similarity between two variables (1).

$$r_{X,Y} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2}}, \quad (1)$$

where: x_i, y_i : individual sample points;

$\bar{x}, \bar{y}$: sample means;

n : number of observations.

A correlation value close to +1 or -1 indicates a strong linear relationship, while a value near 0 suggests no linear correlation. To minimize redundant input information to the model and reduce the dimensionality of the feature space, features with correlation of $|r| > 0.8$ were examined and closely and staged for elimination to the modeling step. Additionally, the corresponding matrix (Fig. 2) helped us eliminate perfectly correlated features that naturally could be foreshadowed by the constraints of the physical system. For example, it's logical that when the yaw rate of the ship increases, its velocity at the stern will also increase during the turn.

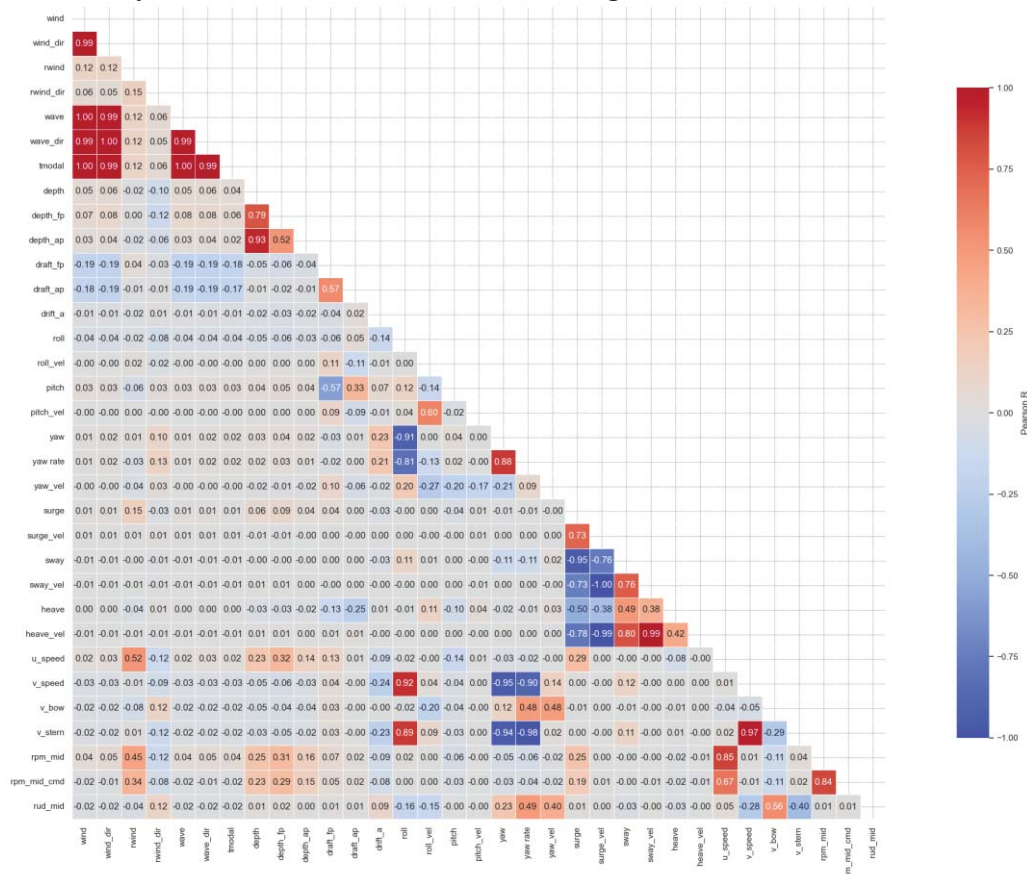


Figure 2 – Pearson R values calculated for features in the dataset

With the intention of conducting a thorough analysis, it's usually necessary to supplement the Pearson Coefficient with another measure of similarity, such as Spearman's rank or Mutual Information [14]. Since Pearson r assumes the underlying relationship is symmetric and homoscedastic, it can miss or underestimate important nonlinear or monotonic patterns. Mutual Information [13] metrics was used as another similarity measure. Overall, it quantifies the total dependency between variables, measures the information one variable contains about another, and captures both linear and nonlinear dependencies (2). By combining Pearson correlation with Mutual Information, one obtains a more holistic understanding of the structure and dependencies in the data [15].

$$I(X; Y) = \sum_{x \in X} \sum_{y \in Y} p(x, y) \log \left(\frac{p(x, y)}{p(x) p(y)} \right); \quad (2)$$

$$p(x, y) = \frac{\text{count}(X=x, Y=y)}{N}; \quad (3)$$

$$p(x) = \frac{\text{count}(X=x)}{N}; \quad (5)$$

$$p(y) = \frac{\text{count}(Y=y)}{N}, \quad (6)$$

where $p(x, y)$: joint probability distribution of X and Y (3);
 $p(x), p(y)$: marginal probability distributions (5, 6).

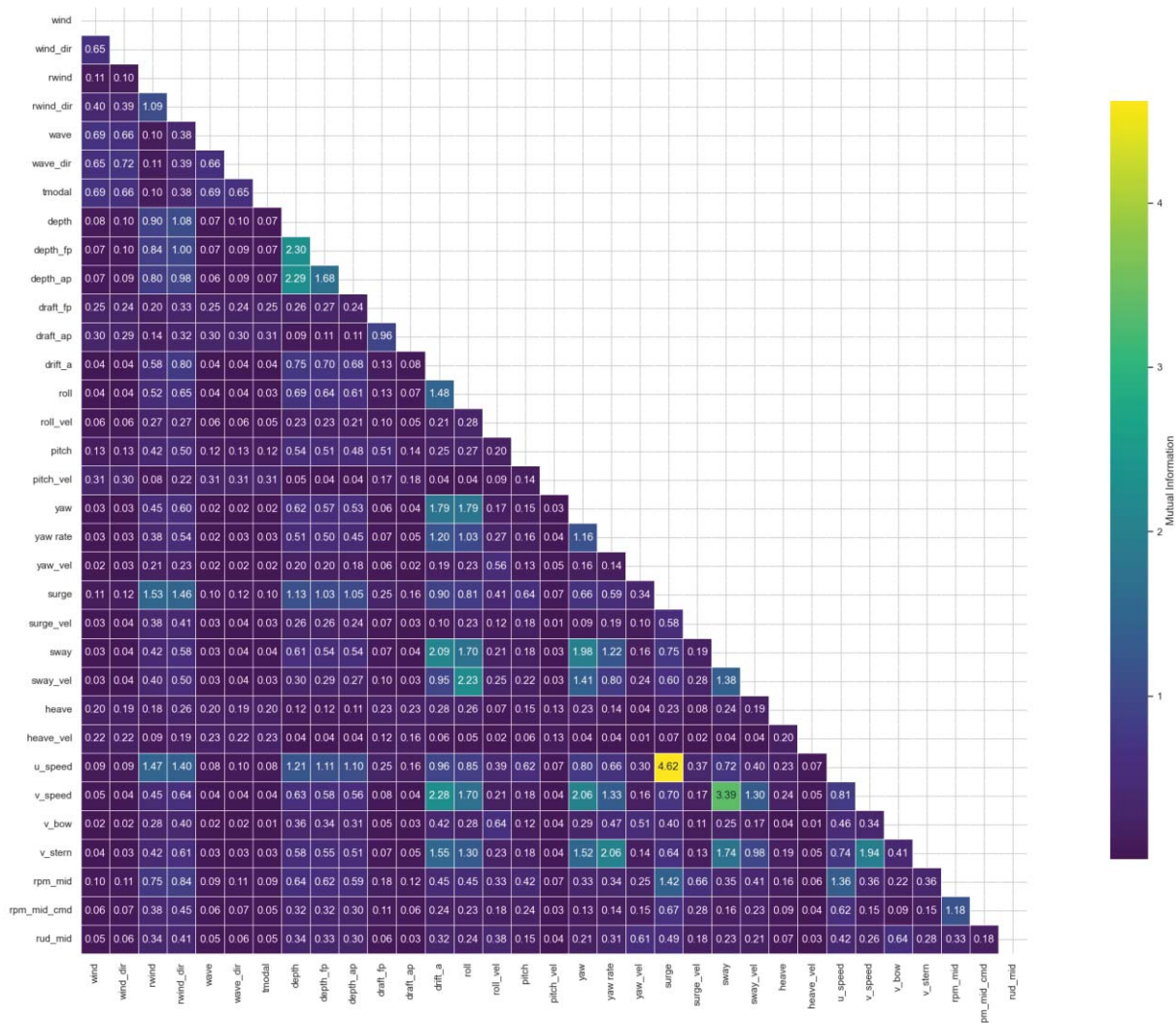


Figure 3 – Mutual Information values calculated for features in the dataset

Applying this technique helped us confirm correlated features from calculating the Pearson Coefficient and identify new naturally outcoming similarities. For example, MI coefficient between *uw_speed* and *u_speed* of 9.257 and *uv_speed* and *v_speed* of 8.065 indicated that given our dataset and for modeling purposes underwater vessel speed wouldn't represent meaningful information.

Following both procedures mentioned above, the list of meaningful features was reduced to a tight list that will provide meaningful information. The features of control such as current mid-engine RPM (*rpm_mid*), commanded mid-engine RPM (*rpm_mid_cmd*), current rudder angle (*rud_mid*), commanded rudder angle (*rud_mid_cmd*). The trajectory or directional features include the eastward position coordinate (*pos_east*) and northward position coordinate (*pos_north*). The motion features describe the ship's movement on the cartesian coordinate system, such as *pitch*, *roll*, *yaw*, *sway*, and *surge*. Velocity features forward surge velocity (*u_speed*), sideways sway velocity (*v_speed*), velocity at the bow (*v_bow*) and velocity at the stern (*v_stern*). While it's acknowledged that weather-related features might trigger different behaviors of navigators, weather-related features were decided to be eliminated since the data in our experiment contained almost identical weather conditions.

Modelling. A novel unsupervised learning pipeline is proposed to facilitate actionable, interpretable feedback for cadets and instructors participating in maritime simulation exercises (Fig. 4). First, the multivariate time series data was preprocessed for each session using established time series feature engineering techniques. This includes calculating aggregate statistics (e.g., mean, standard deviation, minimum, maximum, skewness, and kurtosis) over rolling windows, as well as generating lagged features that capture recent trends and temporal dependencies in the navigational telemetry. The resulting set of fixed-length feature vectors represents each exercise session in a manner that preserves relevant temporal dynamics while remaining interpretable. These feature vectors are then clustered using a Gaussian Mixture Model (GMM), enabling the identification of groups of cadets exhibiting similar navigational behavior or performance patterns. To further enhance the interpretability of each cluster, an XGBoost classifier was trained [17] as a surrogate model to predict cluster membership based on the engineered features. The feature importances and decision logic of this model are subsequently analyzed using SHapley Additive exPlanations (SHAP) [18], providing domain-relevant, actionable insights into the behavioral characteristics defining each group. This integrated approach supports data-driven instructional feedback and personalized training interventions within maritime education. The resulting explanations enable instructors to target training interventions toward specific performance factors, ultimately supporting more effective and individualized learning pathways.



Figure 4 – Modeling pipeline overview

Preprocessing. A systematic feature engineering was performed on the raw navigational telemetry data to enable robust, session-level analysis and clustering. In order to mitigate small sample size of 13 sessions we decided to cluster navigator behavior on events within the session by partitioning each session into fixed-length windows of 10 minutes with 50% overlap, generating roughly 15–20 segments per session. The total number that our method generated was 201. For each segment, a set of global statistical descriptors was extracted from the original features to summarize temporal dynamics while reducing dimensionality.

Specifically, for each numeric telemetry feature, a suite of standard aggregate statistics was calculated – including the mean, standard deviation, minimum, and maximum values – across all timesteps within a session. Mathematically, for a feature x and session s , it was computed:

$$\text{Mean: } \mu_{x,s} = \frac{1}{T} \sum_{t=1}^T x_{s,t}$$

$$\text{Standard deviation: } \sigma_{x,s} = \sqrt{\frac{1}{T} \sum_{t=1}^T (x_{s,t} - \mu_{x,s})^2},$$

Minimum: $\min_{x,s} = \min_t x_{s,t}$,

Maximum: $\max_{x,s} = \max_t x_{s,t}$,

where T denotes the number of timesteps in session s . This procedure was implemented by grouping the data by *series_id* and slicing desired window interval, and applying the aggregation functions to each feature, resulting in a single summary vector per segment. The resulting feature matrix thus captures both central tendencies and variability for each navigational segment, laying the foundation for interpretable clustering and subsequent analysis (see Figure 4 for an overview of the preprocessing workflow).

Clustering. Given our aim to identify natural groupings within rolling-, several clustering algorithms was considered: KMeans [19], Gaussian Mixture Models (GMM) [20], DBSCAN [21], and HDBSCAN [22]. KMeans and GMMs are widely used centroid-based approaches that assume convex, isotropic structures and require the number of clusters k to be specified in advance. This poses significant limitations for our application, as the actual number of navigational behavior types is unknown, and the expected cluster shapes may be non-spherical and of varying density. Moreover, these algorithms are sensitive to noise and outliers, which might be problematic given our small sample size. DBSCAN addresses many of these challenges by defining clusters as dense regions of data points separated by areas of lower point density. However, DBSCAN is limited by its reliance on a global density threshold ε , making it ill-suited to discover clusters of differing densities – a characteristic anticipated in our rolling-window aggregated representations.

To overcome these limitations, HDBSCAN (Hierarchical Density-Based Spatial Clustering of Applications with Noise), was selected, which generalizes DBSCAN by allowing the detection of clusters with variable densities and does not require the number of clusters to be specified. HDBSCAN is thus highly suitable for our experimental context, where data may exhibit clusters of widely varying density and shape and where noise and outliers are prevalent.

HDBSCAN adopts the core distance concept initially introduced in the DBSCAN and LOF literature. For a point x , the core distance concerning a parameter k is denoted as $core_k(x)$ (7) and formally defined as the distance from x to its k -th nearest neighbor (inclusive):

$$core_k(x) = \text{distance to the } k\text{-th nearest neighbor of } x \quad (7)$$

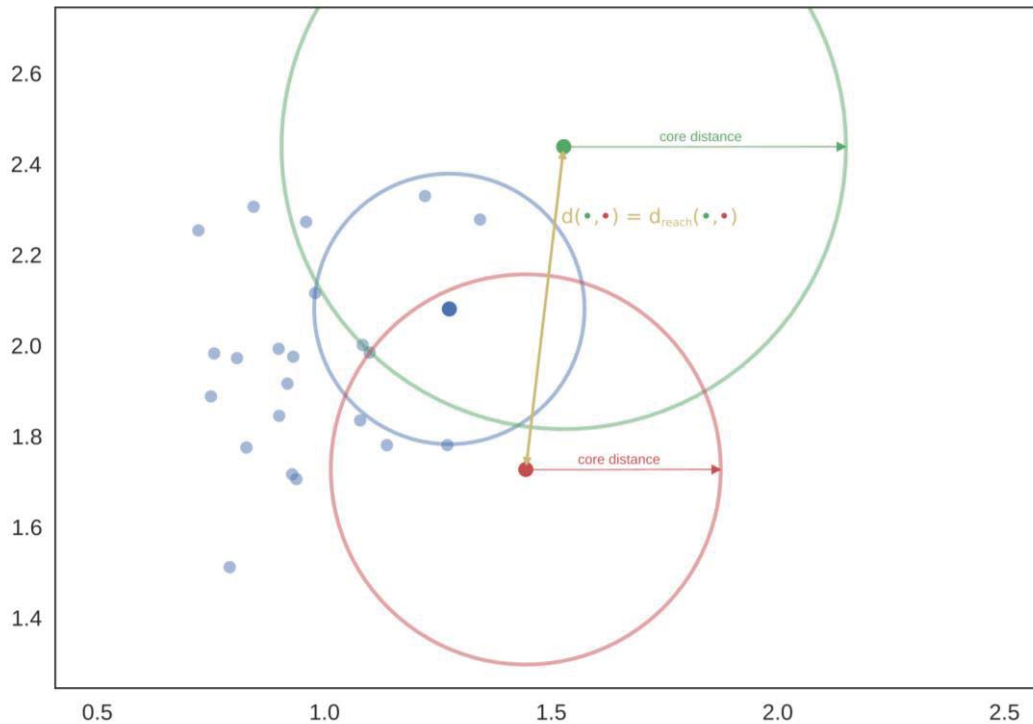


Figure 5 – HDBSCAN reachability distance

However, to effectively distinguish and separate points in sparse areas (with high core distance), HDSCAN introduces the mutual reachability distance (Fig. 5) between two points, a and b , denoted (8)

$$d_{mreach-k}(a, b) = \max\{core_k(a), core_k(b), d(a, b)\}, \quad (8)$$

where $d(a, b)$ is the original metric (e.g., Euclidean) distance between a and b . The multiple different points are depicted with different colors. The next step is to build a Minimum Spanning Tree (MST) using Prim's algorithm – the tree is constructed one edge at a time, always adding the lowest weight edge that connects the current tree to a vertex not yet in the tree (Fig. 6).

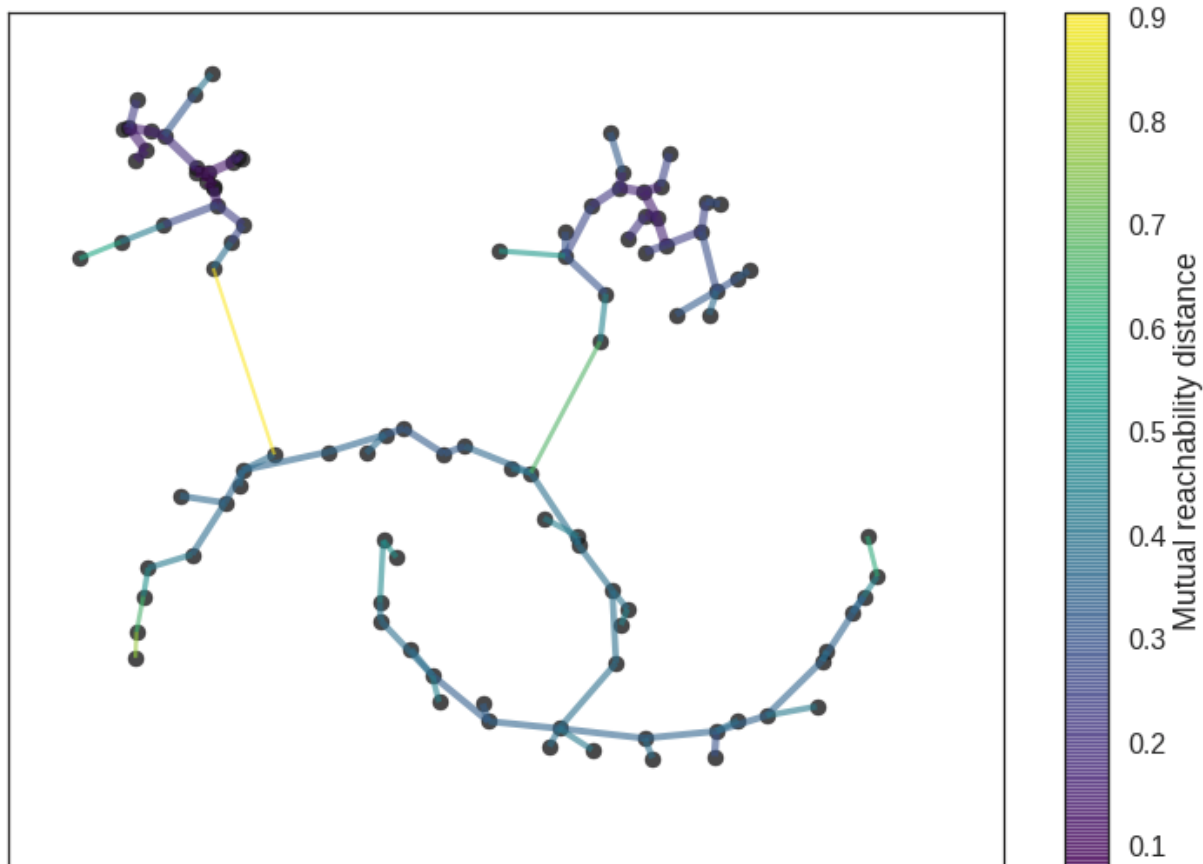


Figure 6 – HDBSCAN's Minimum Spanning Tree (MST)

After obtaining the MST, the conversion to a hierarchy of connected components takes place in reverse order by sorting the edges of the tree by distance and iterating over them using a union-find data structure. Subsequently, the algorithm condenses the hierarchy by eliminating clusters that persist over a small range of distance thresholds and retaining stable clusters over a wide range of thresholds (6). Stability is defined by:

$$\sum_{p \in \text{cluster}} (\lambda_p - \lambda_{\text{birth}}). \quad (9)$$

Finally, HDBSCAN extracts the most stable clusters from the condensed tree. Each data point is assigned to a cluster or noise. Given the performance characteristics, ease of use, and suitability for potential noise simulation data, the DBSCAN algorithm was preferred.

HDBSCAN clustering algorithm was applied to the aggregated session-level feature data, adjusting the *min_cluster_size* parameter to 8 and *min_samples* to 1. t-distributed stochastic neighbors embedding (t-SNE) was employed to project the high-dimensional feature vectors onto the first two components to visualize the clustering structure in a reduced-dimensional space. Figure 7 presents the resulting clusters, with each point representing a navigational segment and colors denoting cluster membership.

Points labeled "-1" correspond to sessions that HDBSCAN could not confidently assign to any cluster, reflecting their outlier or ambiguous nature in the feature space. The visualization reveals that algorithm identified three distinct clusters labeled: "0", "1", and "2". While the amount of noise is noteworthy, the clusters still appear to be distinguishable. The amount of noise can be explained by the inherent discrepancies in real world data which occurs during educational simulation. For example, not all students started performing exercise immediately even though the session was already initiated. Interferences of this kind is inevitable and therefore, it's important to build systems that can understand those patterns and know how to distinguish them.

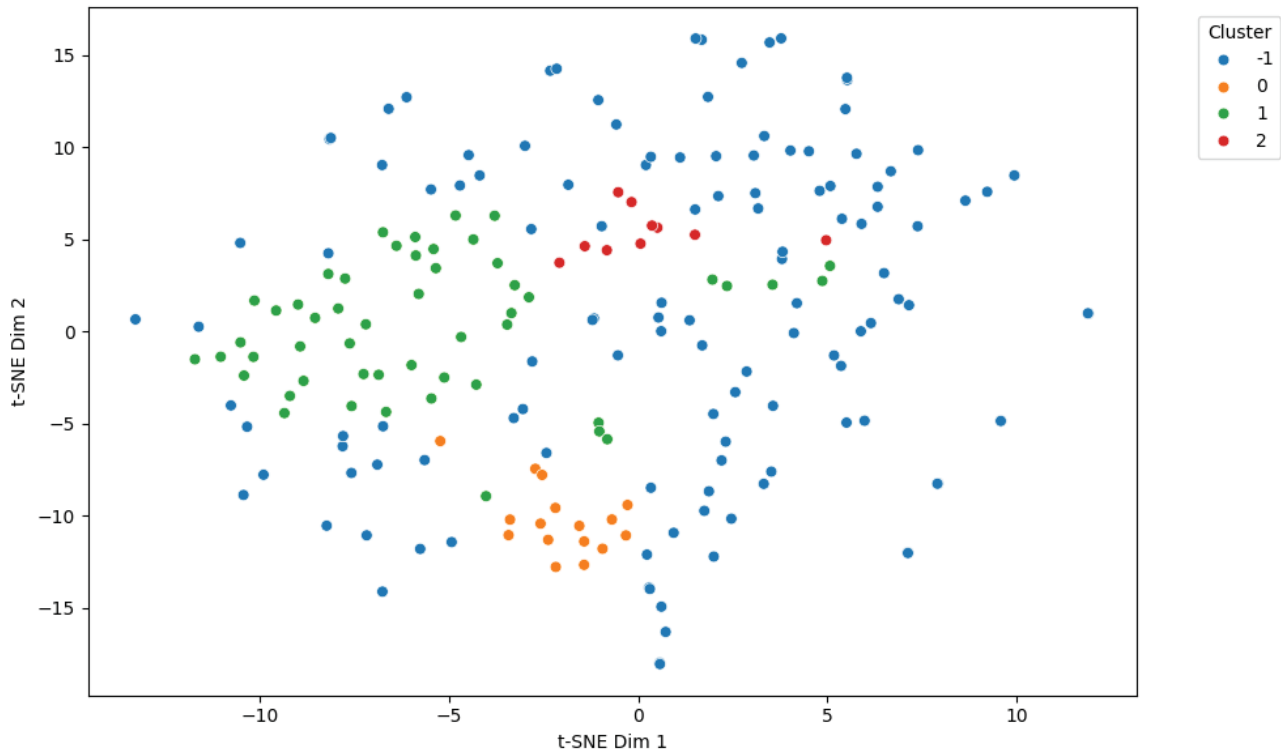


Figure 7 – HDBSCAN clusters visualized using t-SNE

Support Vector Machines SVM classifier was trained to interpret clustering results on the obtained cluster labels. SVMs are supervised learning models well-suited for classification tasks, particularly with small and high-dimensional datasets [23]. SVMs operate by identifying the optimal separating hyperplane in the feature space that maximally distinguishes between different classes. In its linear form, the SVM seeks the hyperplane defined by the equation (10):

$$w^T x + b = 0, \quad (10)$$

where w is the normal vector to the hyperplane and b is the intercept. The decision rule for class assignment is based on the sign of this function. The optimal hyperplane (11) finds the optimal by maximizing the margin – the distance between the hyperplane and the closest data points from each class (called support vectors):

$$\max_{w, b} \left(\frac{2}{|w|} \right); \quad (11)$$

$$\text{subject to } y_i(w^T x_i + b) \geq 1 \quad \forall i,$$

where $y_i \in \{-1, 1\}$ are class labels (Fig. 8).

To handle nonlinear class boundaries, SVMs can employ the kernel trick, implicitly mapping data into a higher-dimensional feature space where linear separation is feasible [24]. However, in our application, a linear SVM was employed for maximum interpretability and computational efficiency, given the limited sample size.

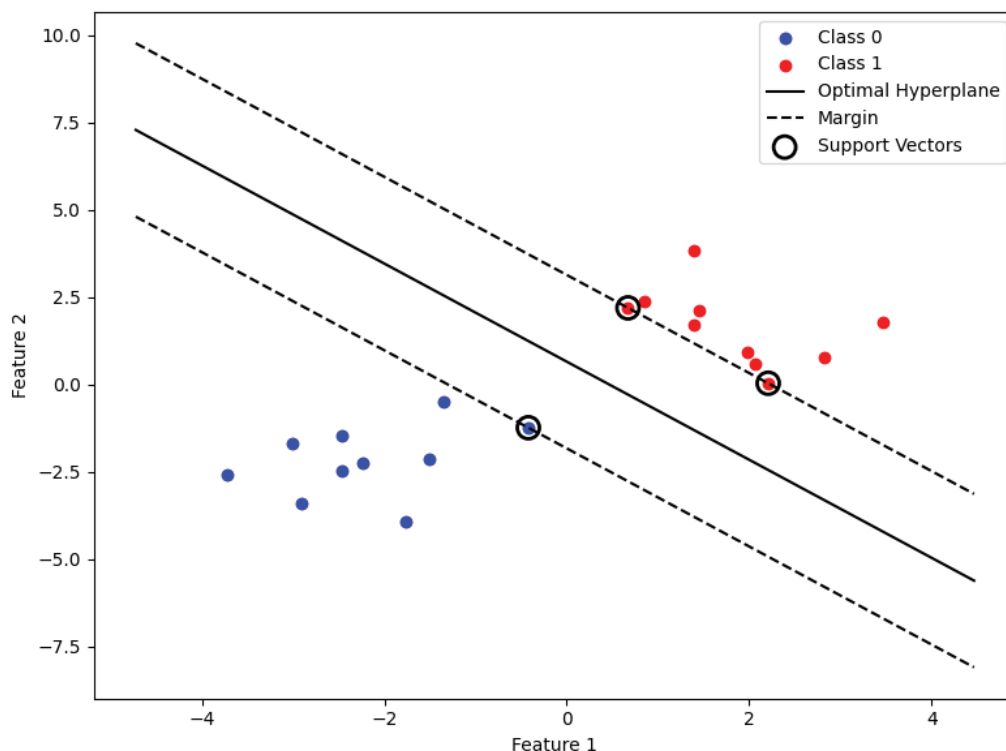


Figure 8 – Illustration of Linear SVM: Support Vectors and Margin

Support Vector Machines were chosen as the surrogate model for cluster explanation due to several compelling advantages in our experimental context. First, SVMs are particularly well-suited for small datasets because the support vectors determine their decision boundary. This helps mitigate the risk of overfitting and enhances robustness to limited sample sizes [25]. Moreover, in the linear case, SVMs are inherently interpretable, as the coefficients of the decision function w directly quantify the importance of each feature in distinguishing between clusters, providing transparent insight into the classification process [26].

SVMs also integrate effectively with SHAP (SHapley Additive exPlanations) [27], allowing us to quantify the contribution of individual features to each cluster assignment and thus facilitate nuanced, quantitative analysis of navigational behaviors, even in multi-class scenarios. Furthermore, SVMs are relatively robust to the curse of dimensionality and do not require extensive hyperparameter tuning, which is particularly advantageous given our small dataset and the engineered feature space. These properties make the SVM classifier an ideal and interpretable surrogate for elucidating which engineered features most influence the assignment of navigational sessions to each cluster, thereby bridging unsupervised clustering with actionable, human-understandable feedback for instructors and cadets.

Interpretability. To translate model predictions into actionable feedback for cadets and instructors, SHAP (SHapley Additive exPlanations) was utilized to attribute each cluster assignment to its underlying feature contributions. SHAP is a state-of-the-art, model-agnostic framework for explainable artificial intelligence [27]. It is grounded in cooperative game theory, specifically the concept of Shapley values, which allocate credit for a model's output among its input features. For each prediction, the SHAP value ϕ_i (12) for feature j is computed as the average marginal contribution of that feature across all possible feature subsets:

$$\phi_j = \sum_{S \subseteq F \setminus \{j\}} \frac{|S|! (|F| - |S| - 1)!}{|F|!} [f_{S \cup \{j\}}(x_{S \cup \{j\}}) - f_S(x_S)], \quad (12)$$

where F is the set of all features, S is a subset of features not containing j , and f_S is the model trained on features in S . This formulation ensures that feature attributions are consistent, locally accurate, and add up to the difference between the model's prediction and its expected value [28].

For each session in our dataset, SHAP values were computed for all features attributing the factors underlying each cluster assignment. These attributions were visualized using SHAP summary and bar plots (Fig. 9), which rank features by their global importance and show both the magnitude and direction of their effect on cluster membership. This interpretability framework identifies the most influential navigational features across all cadets and supports actionable, individualized feedback by revealing which specific behaviors most contributed to a segment's cluster assignment. The feedback can be traced to specific segment in the session and therefore can be traced back to the exact location of occurrence. As a result, SHAP bridges the gap between complex machine learning outputs and human-understandable instructional insights.

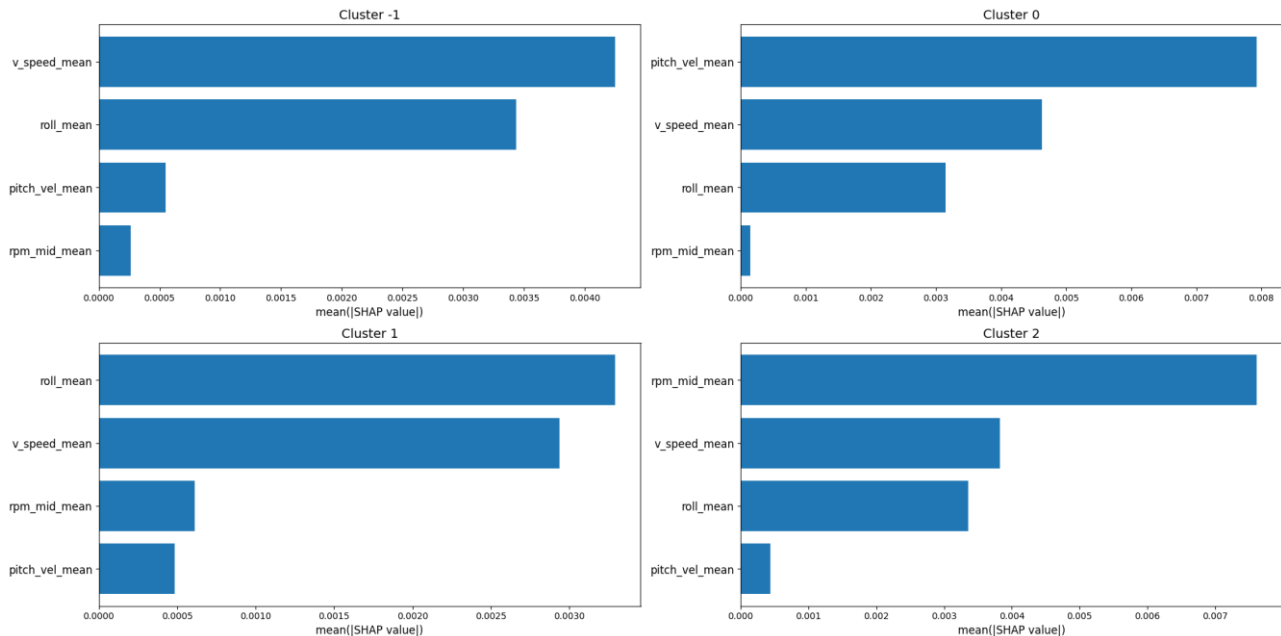


Figure 9 – Illustration of SHAP values for each identified cluster

The SHAP value analysis for each cluster – corresponding to distinct patterns of navigational performance, including noise/outlier sessions – reveals the features most influential in determining cluster assignments.

For sessions identified as outliers or noise (Cluster -1) and for sessions in Cluster 0, the mean vertical speed (*v_speed_mean*), average roll angle (*roll_mean*), followed by average rate of change of pitch (*pitch_vel_mean*) and average mid-propeller RPM (*rpm_mid_mean*) are the top contributors to the cluster attribution. This combination of elevated lateral motion and vessel roll suggests instability – possibly due to corrections and rude steering. The lower value of propeller PRM in the Cluster 0 may indicate attempt to maneuver at lower speeds.

For Cluster 1, the mean roll velocity (*roll_mean*), mean longitudinal speed (*v_speed_mean*), and average propeller RPM (*rpm_mid_mean*) are defining features. This perhaps could indicate stable vessel motion where cadets exhibit stable navigation with consistent lateral motion and propulsion given relatively moderate SHAP values across features and the fact that this corresponds to the largest cluster identified.

Segments in Cluster 2 show high importance of average propeller RPM (*rpm_mid_mean*), suggesting that cadets in this group rely heavily on propulsion power for navigation. This behavior may be required in certain contexts but might also lead to sharp or energy-inefficient maneuvers.

While those suggestions offer global patterns of behavior, the system can trace back the segment classified as certain cluster and provide detailed view on what was happening with the vessel and navigator at that time interval. This should greatly assist and quicken instructor's ability to provide actionable feedback to cadets.

Conclusion. This study has demonstrated the feasibility and utility of applying unsupervised machine learning and explainable AI techniques to maritime simulator data for the automated assessment of cadet navigational performance. By engineering interpretable features from raw time-

series telemetry and leveraging advanced clustering algorithms such as HDBSCAN, the proposed pipeline effectively identifies meaningful patterns in cadet behaviors during simulation exercises. The integration of surrogate models and SHAP-based interpretability provides actionable feedback, enabling instructors to tailor training interventions to address specific performance gaps and reinforce safe navigational strategies. Notably, the results underline the pivotal role of key features—such as roll velocity, propeller RPM, and sideways velocity – in distinguishing between stable, dynamic, and anomalous navigational styles. These findings support the transition from subjective, instructor-driven assessments toward a more objective, data-driven approach in maritime education, with the potential to reduce human error and enhance overall maritime safety.

Prospects for further research. Despite promising results, several avenues remain open for future research. First, expanding the dataset to include a larger and more diverse sample of cadets and navigational scenarios will strengthen the generalizability and robustness of the findings. Future work may also focus on real-time deployment of the feedback system, enabling adaptive, in-situ guidance during simulation or actual vessel operation. Moreover, integrating the proposed method with fuzzy logic-based risk assessment frameworks may further enhance the system's capability to support decision-making under uncertainty. Finally, collaboration with maritime training institutions and stakeholders will be essential for validating the practical impact of automated feedback tools and for driving the evolution of competency-based, individualized maritime education.

REFERENCES

1. Galieriková, A. (2019). The human factor and maritime safety. *Transportation Research Procedia*, 40, 1319–1326. <https://doi.org/10.1016/J.TRPRO.2019.07.183>.
2. International Maritime Organization. (n.d.). *The human element*. IMO. Retrieved May 24, 2025, from <https://www.imo.org/en/OurWork/HumanElement>.
3. Munim, Z. H., Dushenko, M., Jimenez, V. J., Shakil, M. H., & Imset, M. (2020). Big data and artificial intelligence in the maritime industry: a bibliometric review and future research directions. *Maritime Policy & Management*, 47(5), 577–597. <https://doi.org/10.1080/03088839.2020.1788731>.
4. Nguyen, D., Vadaine, R., Hajduch, G., Garelo, R., & Fablet, R. (2021). GeoTrackNet-A Maritime Anomaly Detector using Probabilistic Neural Network Representation of AIS Tracks and A Contrario Detection. *IEEE Transactions on Intelligent Transportation Systems*, 1–13. <https://doi.org/10.1109/TITS.2021.3055614>.
5. Ana, Mateus Sant', Guoyuan Li, and Houxiang Zhang. "A Decentralized Sensor Fusion Approach to Human Fatigue Monitoring in Maritime Operations." *2019 IEEE 15th International Conference on Control and Automation (ICCA)*, July 1, 2019, 1569–74. <https://doi.org/10.1109/ICCA.2019.8899708>.
6. Xiao, Z., Fu, X., Zhang, L., Zhang, W., Liu, R. W., Liu, Z., & Goh, R. S. M. (2020). Big Data Driven Vessel Trajectory and Navigating State Prediction With Adaptive Learning, Motion Modeling, and Particle Filtering Techniques. *IEEE Transactions on Intelligent Transportation Systems*, 1–14. <https://doi.org/10.1109/TITS.2020.3040268>.
7. BIMCO/ICS, 2021. Seafarer Workforce Report: The global supply and demand for seafarers in 2021. Available at: <https://www.ics-shipping.org/publication/seafarer-workforce-report-2021-edition/>.
8. Eurydice (2021). Draft Law on the Organization and Integration of Vocational Training. Available at: https://eacea.ec.europa.eu/national-policies/eurydice/content/national-reforms-vocational-education-and-training-and-adult-learning-70_en.
9. Munim, Ziaul & Kim, Tae-eun. (2023). A Review of Learning Analytics Dashboard and a Novel Application in Maritime Simulator Training. 10.54941/ahfe1003158.
10. W. C. (2025). *Wärtsilä Navigation simulator NTPRO 5000*. Wärtsilä Navigation Simulator NTPRO 5000. Retrieved May 27, 2025, from

<https://www.wartsila.com/marine/products/simulation-and-training/navigational-simulators/navigation-simulator-ntpro-5000>.

11. Good, I. J. The Philosophy of Exploratory Data Analysis. *Philosophy of Science*. 1983;50(2):283–295. <https://doi.org/10.1086/289110>.
12. Benesty, J., Chen, J., Huang, Y., Cohen, I. (2009). Pearson Correlation Coefficient. In: Noise Reduction in Speech Processing. Springer Topics in Signal Processing, vol 2. Springer, Berlin, Heidelberg. https://doi.org/10.1007/978-3-642-00296-0_5.
13. Veyrat-Charvillon, N., & Standaert, F. X. (2009). Mutual information analysis: how, when, and why?. In *International workshop on cryptographic hardware and embedded systems* (pp. 429–443). Berlin, Heidelberg: Springer Berlin Heidelberg.
14. Martínez Vásquez, D. A., Posada-Quintero, H. F., Rivera Pinzón, D. M. (2023). Mutual Information between EDA and EEG in Multiple Cognitive Tasks and Sleep Deprivation Conditions. *Behavioral Sciences*. 2023; 13(9):707. <https://doi.org/10.3390/bs13090707>.
15. Kamalov, F., Sulieman, H., Alzaatreh, A., Emarly, M., Chamlal, H., Safaraliev, M. (2025). Mathematical Methods in Feature Selection: A Review. *Mathematics*. 2025; 13(6):996. <https://doi.org/10.3390/math13060996>.
16. Yong, Yu, Xiaosheng, Si, Changhua, Hu, Jianxun Zhang (2019). A Review of Recurrent Neural Networks: LSTM Cells and Network Architectures. *Neural Comput* 2019; 31 (7): 1235–1270. https://doi.org/10.1162/neco_a_01199.
17. Chen, T., & Guestrin, C. (2016). XGBoost. KDD '16: Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 785–794. <https://doi.org/10.1145/2939672.2939785>.
18. LIME: Ribeiro, Marco Tulio, Sameer Singh, and Carlos Guestrin. "Why should I trust you?: Explaining the predictions of any classifier." Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. ACM, 2016.
19. Ahmed, M., Seraj, R. (2020). Islam SMS. The K-means Algorithm: A Comprehensive Survey and Performance Evaluation. *Electronics*; 9(8):1295. <https://doi.org/10.3390/electronics9081295>.
20. Wan, H., Wang, H., Scotney, B., and Liu, J. (2019). "A Novel Gaussian Mixture Model for Classification," 2019 IEEE International Conference on Systems, Man, and Cybernetics (SMC), Bari, Italy, pp. 3298–3303, <https://doi.org/10.1109/SMC.2019.8914215>.
21. Khan, K., Rehman, S. U., Aziz, K., Fong, S. and Sarasvady, S. (2014). "DBSCAN: Past, present and future," *The Fifth International Conference on the Applications of Digital Information and Web Technologies (ICADIWT 2014)*, Bangalore, India, pp. 232–238, <https://doi.org/10.1109/ICADIWT.2014.6814687>.
22. Campello, R.J.G.B., Moulavi, D., Sander, J. (2013). Density-Based Clustering Based on Hierarchical Density Estimates. In: Pei, J., Tseng, V.S., Cao, L., Motoda, H., Xu, G. (eds) Advances in Knowledge Discovery and Data Mining. PAKDD 2013. Lecture Notes in Computer Science, vol 7819. Springer, Berlin, Heidelberg. https://doi.org/10.1007/978-3-642-37456-2_14.
23. Cortes, C., Vapnik, V. (1995). Support-vector networks. *Mach Learn* 20, 273–297, 1995. <https://doi.org/10.1007/BF00994018>.
24. Cristianini, N., Shawe-Taylor, J. (2000). *An Introduction to Support Vector Machines and Other Kernel-Based Learning Methods*. Cambridge University Press.
25. Hsu, C.-W., Chang, C.-C., & Lin, C.-J. (2016). A Practical Guide to Support Vector Classification.
26. Guyon, I., Weston, J., Barnhill, S., & Vapnik, V. (2002). *Gene Selection for Cancer Classification using Support Vector Machines*. Machine Learning, 46(1), 389–422.
27. Lundberg, S. M., & Lee, S.-I. (2017). A Unified Approach to Interpreting Model Predictions. Advances in Neural Information Processing Systems, 30, 4765–4774.
28. Strumbelj, Erik and Kononenko, Igor (2010). *An Efficient Explanation of Individual Classifications using Game Theory*. J. Mach. Learn. Res. 11 (3/1/2010), 1–18.

Кочубей П., Носов П. ВИЗНАЧЕННЯ ВПЛИВУ ФАКТОРУ ОПЕРАТОРІВ-СУДНОВОДІЇВ ЗАСОБАМИ НАВІГАЦІЙНОГО ТРЕНАЖЕРУ

У міру ускладнення морських операцій зростає потреба у впровадженні об'єктивних та масштабованих методів оцінювання компетентності моряків, що виходять за межі традиційних інструкторських підходів. У цьому дослідженні представлено комплексну, базовану на даних систему аналізу роботи курсантів у тренажерних вправах, яка використовує сучасні методи без наочного навчання та пояснюваного штучного інтелекту (AI). Для аналізу використовувалися багатовимірні часові ряди з навігаційних симуляцій, що охоплюють динаміку судна, дії екіпажу та параметри навколишнього середовища за десятками різних ознак. Було реалізовано ретельний етап попередньої обробки даних, який поєднує агрегування статистичних ознак та усунення надлишкової інформації за допомогою коефіцієнта кореляції Пірсона та взаємної інформації. Це дозволило сформувати компактний, але інформативний набір характеристик, який відображає як керуючі впливи, так і навігаційні стани та рух судна. Кожну сесію симуляції було закодовано у вигляді вектора, що фіксує середні значення та варіативність параметрів протягом виконання вправи. Для кластеризації було обрано алгоритм HDBSCAN, який особливо ефективний для виявлення груп із різною щільністю та автоматично виділяє аномальні випадки, що критично важливо для оцінки підготовки. Знайдені кластери візуалізували за допомогою T-розподіленого вкладення стохастичної близькості (t-SNE), що дозволило інтерпретувати патерни дій курсантів. Для пояснення особливостей кожної групи було навчено лінійну SVM-модель, а метод SHAP допоміг проаналізувати, які саме ознаки впливають на рішення моделі. Основні результати показали, що кластери відповідають різним стилям навігації: стабільні, обережні підходи відрізняються від динамічних чи ризикованих за такими характеристиками, як швидкість крену, кутовий рух (yaw_rate) та оберти двигуна. Сесії, віднесені до аномалій, зазвичай характеризуються різкими маневрами чи нестійким керуванням, що може свідчити про наявність прогалин у навичках. Інтерпретовані SHAP-значення перетворюють складні висновки моделі на зрозумілі для інструкторів рекомендації, даючи можливість адресно працювати з недоліками кожного курсанта. Запропонований підхід може стати прозорою та масштабованою альтернативою суб'єктивному оцінюванню у морській освіті, з реальними перспективами для підвищення безпеки та персоналізації навчання. Система показує високий потенціал інтеграції в практичне середовище підготовки кадрів та подальшого розвитку з розширенням масиву доступних даних.

Ключові слова: водний транспорт; експлуатація засобів транспорту; безпека судноплавства; фактор людини; автоматизація; ризик; інтелектуальні системи; інформаційні панелі аналітики навчання (ІПАН).

© Kochubei P., Nosov P.

Статтю прийнято до редакції 16.06.2025