
Advanced AI OHS Professional Certification (Part II) (Canada)

Effective Use of Machine Learning in OHS Decision Making

Adaptive Learning

Concept: Dynamic adjustment of model parameters in response to new data. Related terms: online learning, incremental training. Explanation: In OHS, adaptive learning enables safety models to refine risk predictions as workplace conditions evolve, improving relevance without full retraining cycles.

Algorithmic Bias

Concept: Systematic error favoring certain outcomes due to data or design. Related terms: fairness, discrimination. Explanation: Bias can skew injury forecasts, leading to inequitable resource allocation; mitigation requires diverse datasets and bias-aware validation.

Anomaly Detection

Concept: Identifying data points that deviate markedly from normal patterns. Related terms: Outlier analysis, novelty detection. Explanation: Detects sudden spikes in near-miss reports, signaling emerging hazards before they cause incidents.

Artificial Neural Network (ANN)

Concept: Computational model inspired by biological neurons. Related terms: Deep learning, multilayer perceptron. Explanation: ANNs model complex interactions among exposure, behavior, and environment to predict occupational injury likelihoods.

Attribute Selection

Concept: Process of choosing relevant features for model input. Related terms: feature engineering, dimensionality reduction. Explanation: Selecting ergonomics, shift length, and equipment age improves model accuracy while reducing overfitting risk.

Baseline Model

Concept: Simple reference model against which advanced models are compared. Related terms: benchmark, naive predictor. Explanation: A logistic regression using historical injury counts serves as a baseline to assess gains from more complex ML approaches.

Batch Learning

Concept: Training a model on a complete dataset at once. Related terms: Offline training, full-dataset learning. Explanation: Useful for periodic OHS risk assessments where data accumulation justifies comprehensive re-training.

Bias-Variance Tradeoff

Concept: Balancing model simplicity and flexibility to minimize total error. Related terms: overfitting,

underfitting. Explanation: In safety prediction, high variance may cause erratic hazard alerts, while high bias can miss subtle risk patterns.

Calibration

Concept: Aligning predicted probabilities with observed frequencies. Related terms: Reliability diagram, probability scaling. Explanation: Well-calibrated OHS models ensure that a 70% risk estimate truly reflects a 70% occurrence rate, supporting credible decision-making.

Classification

Concept: Assigning observations to discrete categories. Related terms: categorical prediction, label assignment. Explanation: Classifies work tasks as "low", "moderate", or "high" injury risk, enabling targeted preventive measures.

Clustering

Concept: Grouping similar data points without predefined labels. Related terms: Unsupervised learning, k-means. Explanation: Segments workers by exposure profiles, revealing hidden sub-populations that may benefit from tailored safety programs.

Confusion Matrix

Concept: Table summarizing true vs. Predicted classifications. Related terms: accuracy, precision, recall. Explanation: Provides OHS analysts with insight into false positive alerts (unnecessary interventions) and false negatives (missed hazards).

Cross-Validation

Concept: Technique for estimating model performance on unseen data. Related terms: K-fold, holdout validation. Explanation: Repeatedly splits OHS datasets to ensure risk predictions generalize across different sites or time periods.

Data Augmentation

Concept: Generating synthetic data to expand training sets. Related terms: oversampling, synthetic minority oversampling technique (SMOTE). Explanation: Balances rare severe injury cases with generated samples, improving model sensitivity to critical events.

Data Governance

Concept: Policies and procedures for managing data quality, security, and compliance. Related terms: data stewardship, privacy. Explanation: Ensures OHS datasets respect worker confidentiality while remaining reliable for machine-learning analysis.

Data Imbalance

Concept: Unequal representation of classes within a dataset. Related terms: class skew, minority class. Explanation: Injury events are often far fewer than non-injury records; imbalance requires resampling or cost-sensitive methods to avoid biased predictions.

Data Integration

Concept: Combining disparate data sources into a unified dataset. Related terms: ETL, data fusion.

Explanation: Merges incident reports, sensor readings, and HR records to enrich predictive features for OHS decision support.

Decision Threshold

Concept: Probability cut-off that determines class assignment. Related terms: operating point, ROC curve.

Explanation: Adjusting the threshold influences the trade-off between alert frequency and missed hazards in safety monitoring systems.

Deep Learning

Concept: Subset of machine learning using multi-layer neural networks. Related terms: Convolutional neural network (CNN), representation learning. **Explanation:** Processes raw video from wearable cameras to automatically detect unsafe postures without manual coding.

Dimensionality Reduction

Concept: Reducing the number of input variables while preserving information. Related terms: Principal component analysis (PCA), feature compression. **Explanation:** Simplifies high-frequency sensor data into core movement patterns for faster OHS risk scoring.

Ensemble Methods

Concept: Combining multiple models to improve predictive performance. Related terms: Random forest, boosting. **Explanation:** Merges decision trees trained on different hazard factors, yielding a more robust injury probability estimate.

Evaluation Metric

Concept: Quantitative measure of model performance. Related terms: F1-score, area under the curve (AUC). **Explanation:** In OHS, the AUC indicates how well a model discriminates between safe and unsafe conditions across all thresholds.

Explainable AI (XAI)

Concept: Techniques that make model decisions understandable to humans. Related terms: model interpretability, SHAP values. **Explanation:** Provides safety managers with clear reasons—such as high vibration exposure—behind a predicted injury risk, supporting transparent actions.

Feature Engineering

Concept: Creating informative variables from raw data. Related terms: derived features, transformation.

Explanation: Converts timestamped shift logs into fatigue indices, enhancing the model's ability to predict fatigue-related incidents.

Feature Importance

Concept: Ranking of input variables based on their impact on predictions. Related terms: variable contribution, permutation importance. **Explanation:** Highlights that ergonomic score and PPE compliance are top drivers of predicted injury rates, guiding resource allocation.

Gaussian Process

Concept: Probabilistic model providing predictions with uncertainty estimates. Related terms: kernel

methods, Bayesian regression. Explanation: Offers OHS analysts confidence intervals around risk forecasts, useful for risk-based decision thresholds.

Generalization

Concept: Model's ability to perform well on unseen data. Related terms: out-of-sample performance, transferability. Explanation: A well-generalized safety model maintains accuracy when applied to new facilities or evolving work practices.

Gradient Descent

Concept: Optimization algorithm that iteratively reduces error. Related terms: learning rate, stochastic gradient descent. Explanation: Trains OHS predictive models by minimizing the difference between predicted and actual injury counts across epochs.

Hyperparameter Tuning

Concept: Adjusting non-learnable settings to improve model performance. Related terms: Grid search, random search. Explanation: Finds the optimal number of trees in a random forest that balances OHS prediction accuracy with computational cost.

Imbalanced Classification

Concept: Classification tasks where one class dominates. Related terms: cost-sensitive learning, resampling. Explanation: In injury prediction, the "no-injury" class overwhelms the "injury" class, requiring specialized strategies to avoid trivial models.

Inference

Concept: Applying a trained model to new data to generate predictions. Related terms: prediction, scoring. Explanation: Real-time inference on sensor streams can trigger immediate alerts when unsafe conditions are detected.

Instance-Based Learning

Concept: Learning algorithms that compare new cases to stored examples. Related terms: K-nearest neighbors (k-NN), memory-based. Explanation: Matches a current work scenario to past incidents to suggest similar preventive actions.

Interpretability

Concept: Degree to which a human can understand the cause of a model's output. Related terms: transparency, XAI. Explanation: Essential for OHS stakeholders to trust automated risk scores and act on them responsibly.

Kernel Trick

Concept: Technique that enables linear algorithms to learn non-linear patterns. Related terms: Support vector machine (SVM), RBF kernel. Explanation: Allows OHS models to capture complex relationships between temperature, humidity, and slip-related injuries.

K-Means Clustering

Concept: Partitioning data into K clusters based on distance to centroids. Related terms: Unsupervised

learning, centroid. Explanation: Groups job tasks by similarity in exposure metrics, revealing hidden high-risk clusters for targeted audits.

Label Encoding

Concept: Converting categorical variables into numeric form. Related terms: One-hot encoding, ordinal encoding. Explanation: Transforms equipment type labels into numbers for inclusion in regression models predicting equipment-related incidents.

Learning Curve

Concept: Plot showing model performance versus training data size. Related terms: sample efficiency, convergence. Explanation: Indicates whether additional OHS data will meaningfully improve risk prediction or if model capacity is saturated.

Logistic Regression

Concept: Statistical model for binary outcome prediction. Related terms: logit, probability modeling. Explanation: Provides a baseline OHS model estimating the odds of injury based on factors such as hours worked and safety training completion.

Loss Function

Concept: Metric that quantifies error during model training. Related terms: cost function, optimization target. Explanation: Mean squared error penalizes large deviations between predicted and actual injury counts, guiding the learning process.

Machine Learning Pipeline

Concept: Structured sequence of steps from data ingestion to model deployment. Related terms: workflow, orchestration. Explanation: In OHS, the pipeline may include data cleaning, feature extraction, model training, validation, and automated risk reporting.

Model Drift

Concept: Degradation of model performance over time due to changing data patterns. Related terms: concept drift, retraining. Explanation: As new safety protocols are introduced, a previously accurate OHS model may misclassify hazards unless monitored and updated.

Model Registry

Concept: Centralized catalog of trained models and their metadata. Related terms: Versioning, artifact store. Explanation: Allows OHS teams to track which risk model was used for a specific audit, supporting reproducibility and compliance.

Multiclass Classification

Concept: Predicting one of three or more categories. Related terms: softmax, one-vs-rest. Explanation: Classifies incidents into "minor", "major", and "critical" severity levels, enabling differentiated response plans.

Neural Architecture Search

Concept: Automated method for designing optimal neural network structures. Related terms: AutoML, hyperparameter optimization. Explanation: Finds the most effective deep-learning layout for processing

high-frequency vibration data in equipment-safety monitoring.

Normalization

Concept: Scaling numeric features to a common range. Related terms: standardization, min-max scaling.

Explanation: Ensures that exposure intensity and shift length contribute proportionally to OHS model training.

One-Hot Encoding

Concept: Representing categorical variables as binary vectors. Related terms: Dummy variables, indicator matrix.

Explanation: Encodes shift type (day, night, swing) for inclusion in a logistic regression predicting fatigue-related injuries.

Outlier Removal

Concept: Excluding extreme values that may distort model training. Related terms: noise filtering, data cleaning.

Explanation: Removes erroneous sensor spikes that could falsely indicate hazardous conditions in safety dashboards.

Overfitting

Concept: Model captures noise rather than underlying pattern, performing poorly on new data. Related terms: high variance, regularization. Explanation: An overfit OHS model may flag routine tasks as high risk, leading to alert fatigue among supervisors.

Parameter

Concept: Learnable coefficient within a model. Related terms: Weight, bias term. Explanation: In a linear OHS risk model, each predictor (e.G., Exposure level) has an associated weight indicating its influence on injury probability.

Precision

Concept: Proportion of positive predictions that are correct. Related terms: positive predictive value, accuracy. Explanation: High precision in safety alerts means most warnings correspond to genuine hazards, preserving trust in the system.

Probabilistic Forecasting

Concept: Generating predictions expressed as probability distributions. Related terms: uncertainty quantification, predictive intervals. Explanation: Provides OHS managers with a range (e.G., 10-15% Chance of a slip) rather than a single deterministic value, supporting risk-based decision making.

Principal Component Analysis (PCA)

Concept: Technique that transforms correlated variables into uncorrelated components. Related terms: eigenvectors, dimensionality reduction. Explanation: Condenses dozens of sensor metrics into a few principal components that capture most variance in workplace motion patterns.

Probabilistic Graphical Model

Concept: Representation of variables and their conditional dependencies using graphs. Related terms:

Bayesian network, Markov model. Explanation: Models causal pathways from equipment failure to injury,

enabling inference about the most likely root causes.

Recall

Concept: Proportion of actual positives correctly identified. Related terms: sensitivity, true positive rate.

Explanation: In OHS, high recall ensures most true hazards trigger alerts, reducing missed-incident risk.

Regularization

Concept: Adding a penalty to the loss function to discourage complex models. Related terms: L1, L2, ridge.

Explanation: Prevents OHS models from over-reacting to noisy data, yielding smoother risk scores.

Reinforcement Learning

Concept: Learning optimal actions through trial-and-error interaction with an environment. Related terms: policy, reward function. Explanation: Can simulate safety interventions, rewarding actions that reduce predicted injury probability, to suggest optimal preventive strategies.

Resampling

Concept: Adjusting the training dataset by oversampling minority class or undersampling majority class.

Related terms: SMOTE, bootstrap. Explanation: Balances scarce severe injury records with duplicated or synthetic examples, improving OHS model sensitivity to critical events.

Risk Matrix

Concept: Visual tool that plots likelihood against severity. Related terms: heat map, risk assessment.

Explanation: Machine-learning outputs can populate a dynamic risk matrix, automatically updating positions as new data arrives.

ROC Curve

Concept: Plot of true positive rate versus false positive rate at various thresholds. Related terms: AUC, diagnostic ability. Explanation: Evaluates OHS model's discriminative power; a curve closer to the top-left corner indicates better hazard detection.

Sampling Bias

Concept: Systematic error introduced when the sample is not representative of the population. Related terms: selection bias, coverage error. Explanation: If only high-visibility incidents are recorded, the model may underestimate hidden low-visibility risks.

Scalable Architecture

Concept: System design that can handle growth in data volume and processing demands. Related terms: distributed computing, cloud deployment. Explanation: Enables OHS predictive services to expand from a single plant to a national network without performance loss.

Segmentation

Concept: Dividing a dataset into distinct groups for targeted analysis. Related terms: Clustering, stratification. Explanation: Segments workers by age and experience to tailor safety messaging based on model-identified risk differentials.

Shapley Additive Explanations (SHAP)

Concept: Method that attributes contribution of each feature to a specific prediction. Related terms: feature attribution, explainability. Explanation: Shows that high noise exposure contributed 0.25 To a worker's predicted injury risk, supporting transparent communication.

Signal-to-Noise Ratio (SNR)

Concept: Measure of useful information relative to background variability. Related terms: data quality, filtering. Explanation: High SNR in vibration sensor data improves the reliability of ML-derived fatigue risk scores.

Simulation Modeling

Concept: Creating virtual representations of real-world processes to test scenarios. Related terms: digital twin, Monte Carlo. Explanation: Generates synthetic incident streams to stress-test OHS predictive models before deployment.

Smart Sensor

Concept: Device that captures data and performs edge analytics. Related terms: IoT, embedded AI. Explanation: Sends real-time temperature and humidity readings to an OHS risk model, enabling instantaneous heat-stress alerts.

Softmax Function

Concept: Normalizes a vector of real numbers into probabilities that sum to one. Related terms: multiclass output, activation. Explanation: Used in OHS neural networks to output probabilities for each incident severity class.

Spatial Analysis

Concept: Examination of geographic or layout-based patterns. Related terms: heat mapping, GIS. Explanation: Identifies zones within a facility where injury predictions consistently exceed thresholds, guiding layout redesign.

Specificity

Concept: Proportion of true negatives correctly identified. Related terms: true negative rate, precision (negative). Explanation: High specificity in OHS alerts reduces unnecessary investigations of safe areas, conserving resources.

Stochastic Gradient Descent (SGD)

Concept: Variant of gradient descent that updates parameters using random mini-batches. Related terms: online learning, optimization. Explanation: Enables rapid updating of OHS models as new incident data streams in, without full retraining.

Supervised Learning

Concept: Model training using labeled input-output pairs. Related terms: classification, regression. Explanation: Uses historical injury records (labels) alongside exposure metrics (features) to teach the algorithm to predict future occurrences.

Support Vector Machine (SVM)

Concept: Algorithm that finds the hyperplane maximizing margin between classes. Related terms: kernel, linear classifier. Explanation: Separates high-risk from low-risk work periods based on multivariate sensor inputs, providing clear decision boundaries.

Temporal Drift

Concept: Shift in data patterns over time due to evolving processes or policies. Related terms: concept drift, seasonality. Explanation: After implementing a new lock-out procedure, injury patterns change, requiring the OHS model to adapt to the new temporal distribution.

Transfer Learning

Concept: Reusing a model trained on one task for a related task. Related terms: fine-tuning, domain adaptation. Explanation: Applies a model trained on manufacturing ergonomics to a construction site, reducing data collection burden while retaining predictive power.

Underfitting

Concept: Model is too simple to capture underlying structure, leading to poor performance on both training and test data. Related terms: high bias, low variance. Explanation: A linear OHS model may miss nonlinear interactions between temperature and PPE compliance, resulting in inaccurate risk estimates.

Unsupervised Learning

Concept: Learning patterns from data without explicit labels. Related terms: clustering, dimensionality reduction. Explanation: Discovers hidden groupings of tasks with similar hazard profiles, informing proactive safety inspections.

Validation Set

Concept: Subset of data used to tune model hyperparameters. Related terms: holdout, cross-validation. Explanation: Provides an unbiased estimate of OHS model performance before final testing on unseen data.

Variance

Concept: Measure of model sensitivity to fluctuations in the training data. Related terms: model instability, overfitting. Explanation: High variance OHS models may produce wildly different risk scores when trained on different weeks of incident data.

Weighted Loss

Concept: Assigning different penalties to errors based on class importance. Related terms: cost-sensitive learning, class weighting. Explanation: In injury prediction, false negatives may be penalized more heavily than false positives to prioritize safety.

Zero-Inflated Model

Concept: Statistical model handling excess zeros in count data. Related terms: hurdle model, overdispersion. Explanation: Addresses the many days with zero recorded injuries while still modeling the occasional high-count events in OHS datasets.