



*An Online CPD Course
brought to you by
CEDengineering.ca*

AI-Assisted Environmental Remediation

Course No: C01-043

Credit: 1 PDH

Brian Lisiewski, P.E.



Continuing Education and Development, Inc.

P: (877) 322-5800

info@cedengineering.ca

Table of Contents

Module 1: Introduction — AI as Advisory, Engineer in Charge, and Public Welfare	3
1.1 Public Welfare & Remediation Goals	3
1.2 Engineer’s Oversight and Ethics	3
1.3 Quality Management & Governance Context	4
Module 2: Remediation & Risk Fundamentals — Processes and Data Landscape	5
2.1 Site Characterization and CSM	5
2.2 Remediation and Control Strategies.....	5
2.3 Risk Assessment Basics	6
2.4 Data Challenges & AI Motivation	6
Module 3: Data-Driven Site Characterization — ML for Delineation & Anomaly Detection	7
3.1 ML for Contamination Delineation.....	7
3.2 Anomaly Detection in Monitoring Data.....	7
3.3 Data Fusion for Site Characterization	8
3.4 Limitations and Confirmations	8
Module 4: Predictive Plume Modeling & Scenario Analysis	9
4.1 ML Surrogate Models for Contaminant Transport	9
4.2 Scenario Analysis & Risk Forecasting.....	9
4.3 Human-in-the-Loop for Adaptive Control	10
4.4 Ensuring Safe and Compliant Outcomes	10
Module 5: Risk Prioritization & Decision-Making Metrics.....	11
5.1 Risk Metrics & ML Inputs	11
5.2 Prioritizing Actions or Sites	11
5.3 Handling Uncertainty in Risk Decisions	12
5.4 Documentation & Regulatory Acceptance.....	12
Module 6: Governance, Model Validation & Cybersecurity	13
6.1 Model Verification & Validation (V&V).....	13
6.2 Configuration Control & Change Management	13
6.3 Data Security & Integrity	14
6.4 Professional Defensibility & Communication	14

References15

Module 1: Introduction — AI as Advisory, Engineer in Charge, and Public Welfare

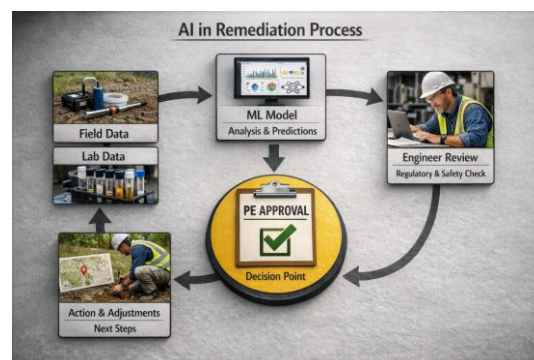
Overview: Cleaning up contaminated sites is fundamentally about protecting people and ecosystems. In this module, we establish how AI and ML can empower environmental professionals to manage contamination risks more proactively, yet never bypass the oversight of a licensed engineer. We frame AI as a tool to enhance site analysis and decision-making — not an autonomous decider. Public welfare (both human health and ecological health) is reinforced: any AI recommendation must be vetted to ensure it does not compromise safety or compliance. The Professional Engineer’s accountability under law and ethics remains absolute, mirroring NSPE guidance that technology does not diminish the engineer’s obligation to safety.

1.1 Public Welfare & Remediation Goals

Environmental remediation projects aim to reduce risk to acceptable levels, whether that means cleaning contaminated soil/groundwater to below certain limits or containing it to prevent exposure. AI can help by identifying hidden contamination patterns or predicting future risks, but all actions derived from AI must ensure no harm is done. For example, if an ML model suggests a less costly remedial plan that leaves contamination in place longer, the PE must verify that this still meets all risk criteria (ensuring no receptors are exposed above guidelines) before considering it — cost savings cannot override safety. Conversely, AI might detect subtle risk pathways (say, a seasonally rising vapor intrusion risk) that an engineer can then address proactively, enhancing public protection.

1.2 Engineer’s Oversight and Ethics

Licensed engineers must approve all site decisions — from delineating “site closure achieved” to selecting remedy options — even if AI aided the analysis. If an ML model classifies a certain area as low-risk and suggests fewer samples, the PE reviews underlying data to confirm no important factor was missed, then decides. This aligns with professional practice standards: the engineer signs off on technical documents (like risk assessments or remediation designs) and thereby attests they have exercised judgment on everything presented, including any AI-derived input. If an ML model indicates an area is likely clean but the engineer’s judgment or field reality says otherwise (e.g., unusual geology or local knowledge suggests a potential hot spot), the engineer’s judgment prevails. This approach ensures AI does not override field experience or regulatory caution.



1.3 Quality Management & Governance Context

Introducing AI into environmental work triggers management and governance responsibilities. We treat an AI model similarly to any critical piece of analysis software: controlled, validated, and documented. This aligns with ISO 9001 principles — e.g., if an ML algorithm is used to estimate contaminant extents, its development, testing, and usage should be documented in the project quality plan, and any updates should go through change control with appropriate sign-offs. Should a regulatory agency ask “How did you determine contamination ends here?”, we need clear evidence (charts, logs, QA tests of the model) showing that determination was based on a sound method under engineer oversight. In essence, AI does not diminish the need for rigorous documentation — it increases it.

Module 2: Remediation & Risk Fundamentals — Processes and Data Landscape

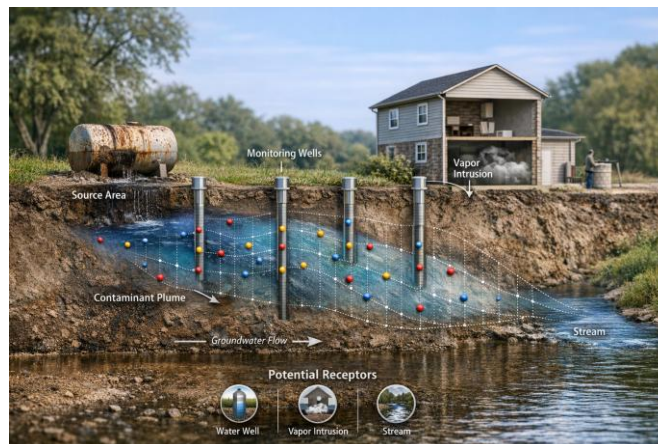
Overview: This module covers the conceptual site model (CSM) approach, typical data collected (soil and groundwater concentrations, flow data, geologic logs), and how risk is quantified (exposure pathways, human/ecological receptors, risk thresholds). We discuss the data landscape — often incomplete and uncertain — and challenges like multiple contaminants or complex hydrogeology, setting the stage to show why advanced analytics like ML can add value by making sense of extensive data and scenarios that are difficult to handle with manual methods alone.

2.1 Site Characterization and CSM

To clean a site, we first build a Conceptual Site Model (CSM) summarizing sources of contamination, how contaminants move (soil, groundwater plumes, vapor migration), and who or what could be exposed. Data underpinning this includes soil boring results (contaminant concentrations at various depths), groundwater monitoring well readings over time, aquifer tests (hydraulic conductivity), and soil properties. In many cases this data is sparse or irregular — e.g., only a handful of well points define a plume of variable shape. AI can assist by learning patterns from limited data plus any proxy data (like geophysical surveys or geochemical indicators) to infer likely contamination extent or hotspot areas that might otherwise be missed. Engineers must cross-check ML outputs with their CSM understanding to ensure algorithmic inferences align with known geologic and chemical principles.

2.2 Remediation and Control Strategies

Remedial approaches (pump-and-treat, soil vapor extraction, in-situ bioremediation, capping, etc.) are chosen based on CSM, site constraints, and risk reduction targets. AI can expedite exploring many such designs by surrogate modeling (see Module 4). However, core constraints remain: the remedy must achieve cleanup goals (like meeting drinking water standards at compliance wells) within a regulatory timeframe or consent order deadline. AI cannot change those goals — it only helps find better ways to meet them. The engineer still ensures any “optimized” plan fulfills regulatory requirements before adoption.



2.3 Risk Assessment Basics

Environmental risk assessment involves assessing exposure pathways (how contaminants could reach humans or ecology) and toxicity. For humans, a risk analysis might calculate a hazard quotient or cancer risk based on pollutant concentration at a receptor and regulatory risk factors. AI can help incorporate more data or simulate many scenarios (Monte Carlo style) for risk — useful when dealing with uncertain plume dynamics or remedy outcomes. Key risk metrics remain determined by regulations and toxicology; AI’s role is to refine predictions or identify which factors drive variance. The engineer communicates risk in time-tested ways, using AI to strengthen confidence in those numbers — not to produce black-box outputs.

2.4 Data Challenges & AI Motivation

Remediation projects face data gaps, uncertainties (e.g., unknown source mass, heterogeneity in geology), and evolving conditions (seasonal changes, new findings). These challenges motivate AI usage:

- For limited monitoring data, ML can extrapolate likely contamination extents or highlight “suspect zones” needing more sampling by learning from analogous patterns or exploring geospatial patterns from the given data.
- With uncertain parameters (like biodegradation rates), ML surrogates can simulate a range of outcomes quickly, helping quantify risk boundaries faster than running full simulations manually — while always referencing back to physical plausibility.
- In complex, multi-contaminant situations, AI clustering can help group co-occurring contaminants to identify shared sources or transport pathways, useful for forensic analysis at sites with multiple contributing industries.

Module 3: Data-Driven Site Characterization — ML for Delineation & Anomaly Detection

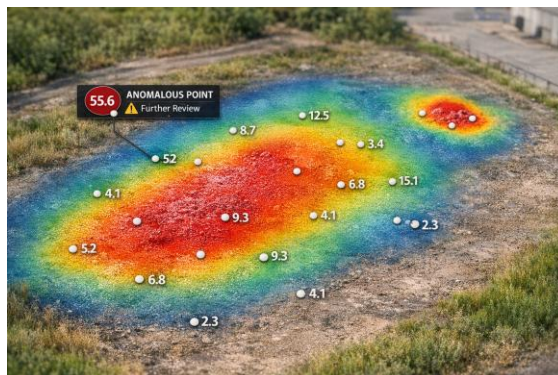
Overview: After collecting site and monitoring data, ML can assist in analyzing these data to build a more complete picture of contamination. This module details how ML models can delineate contamination plumes beyond discrete data points, identify spatial anomalies or outliers that might indicate sampling errors or previously unknown hotspots, and help refine the conceptual site model.

3.1 ML for Contamination Delineation

In site characterization, a common task is estimating contaminant concentration distribution in areas without direct measurements. Traditional approach: Spatial interpolation (like kriging) with domain expert adjustments. ML approach: Training a model (e.g., a support vector machine or neural network) on known data points plus correlated inputs (soil type, depth, distance from source) to predict concentrations at unsampled points. For example, ML surrogate models have been shown to predict key plume features (maximum plume length, spread time, etc.) with approximately 85–90% accuracy by learning from multiple simulated scenarios.¹ Such models can rapidly simulate many contamination configurations under different conditions. The PE must verify that ML predictions align with hydrogeological understanding before relying on them for decisions. If an ML model suggests contamination extends beyond a highway but groundwater flows away from it, the engineer recognizes a possible false positive and may still sample beyond the highway to be safe, rather than relying solely on the ML’s prediction.

3.2 Anomaly Detection in Monitoring Data

AI-based anomaly detection can judge whether a monitoring anomaly is likely an actual event or an outlier (instrument error, lab contamination). For instance, if ML monitors trends and spatial correlation, it might catch that one well’s reading spiked while neighboring wells remained low, flagging it as suspect and recommending a re-sample. Similarly, ML can detect drift in sensor signals over time — if a pH probe gradually reads lower across all wells, the system can alert maintenance because uniform drift suggests a calibration issue, not actual acidification. Professional oversight is essential: the engineer ensures critical decisions are not made automatically based on flagged anomalies, but rather with careful review and verification via duplicate testing before concluding a reading was false.



3.3 Data Fusion for Site Characterization

Often, we have multiple data types: chemical concentrations, electrical resistivity tomography (ERT) indicating conductive plumes, and sentinel vegetation stress indicators. ML can fuse these — e.g., feeding geophysical profiles and chemical data to a model that outputs more robust plume extents than using chemistry alone. Research has shown that sensor data fusion with ML can improve hydrocarbon plume detection reliability compared to single-sensor analysis.¹ The engineer must calibrate and weigh each data source’s trustworthiness and should investigate apparent discrepancies by additional sampling rather than blindly extending the cleanup plan based on model output alone.

3.4 Limitations and Confirmations

While ML expedites site characterization, it does not eliminate the need for ground-truth verification. Regulatory guidance often requires a certain density of sampling to declare an area clean; a model interpolation cannot substitute for that. However, ML can optimize where to sample by highlighting uncertain or critical areas (like plume edges or potential separate plume zones). Any model-introduced biases must be recognized and mitigated — possibly by ensuring predictions are conservative or by applying a safety buffer when interpreting results. ML-assisted delineation must be validated with at least a subset of actual measurements before it is fully trusted for decision-making.

Module 4: Predictive Plume Modeling & Scenario Analysis

Overview: One of AI's most powerful contributions is speeding up “what-if” analyses for contaminant transport. This module explains how ML-based predictive models (surrogates of physics-based groundwater/air models) can simulate contaminant plume evolution and remedial scenarios quickly, letting engineers explore many possibilities in risk assessments or design optimizations. We stress that any such model must be thoroughly validated against trusted numerical models and used with caution.

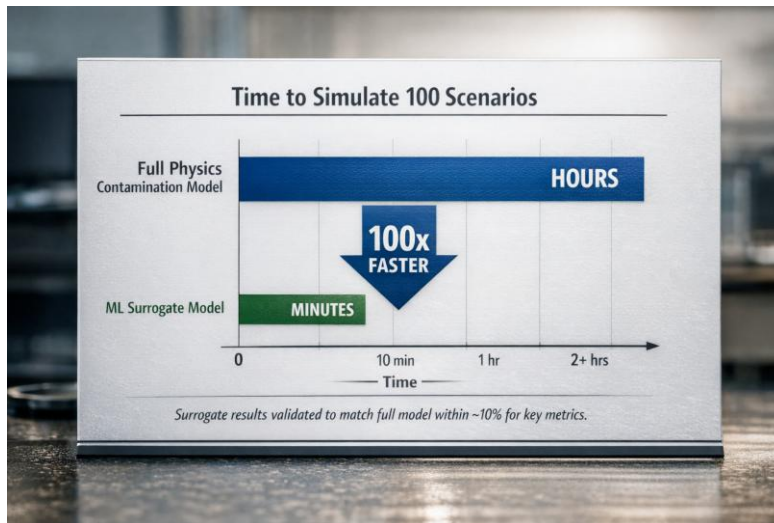
4.1 ML Surrogate Models for Contaminant Transport

Detailed numerical transport models (like finite-difference groundwater solvers) are accurate but slow for repeated runs. ML surrogates can approximate these models effectively. For example, researchers have developed a PSO-SVM surrogate to emulate a 3D groundwater model's key outputs (like plume length and mean concentration), achieving approximately 85–90% agreement with the full numerical model.¹ Another advanced approach integrated a classic hydrogeologic equation (the Thiem well formula) into a deep learning (U-Net) architecture to model dynamic pump-and-treat plume changes, greatly speeding up scenario evaluations at a complex site.² What might otherwise take hours on conventional hardware could be reduced to seconds or minutes with an ML surrogate, depending on model complexity and system configuration — enabling engineers to test dozens of configurations rapidly. The PE must ensure surrogates remain within their valid input domain and double-check a sample of outputs against the original model. Surrogate accuracy may degrade outside trained conditions, so the PE must impose input boundaries and revert to the full physics model if recommendations push those limits.

4.2 Scenario Analysis & Risk Forecasting

AI-driven simulation speed-ups allow robust scenario analysis:

- **Remediation Strategy Comparison:** Evaluate many variations of pumping schedules or injection dosages to see which achieves target concentration fastest. ML can quickly rank these, then the best candidates are run through the full model for confirmation. This ensures the engineer does not miss novel combinations while keeping final picks vetted by high-fidelity simulation and engineering judgment.
- **Natural Attenuation Monitoring:** If relying on monitored natural attenuation (MNA), ML can forecast future plume contraction (or expansion if sources remain) under various attenuation rates, supporting decisions on whether active remediation is needed. These forecasts must align with known hydro geochemistry, which the engineer verifies.
- **Risk Scenario Monte Carlo:** Embedding an ML surrogate in a Monte Carlo simulation varying uncertain parameters thousands of times reveals the distribution of risk outcomes (like how often a well might exceed a threshold). This informs risk management by showing worst-case probabilities — something rarely done traditionally due to computational burden.



4.3 Human-in-the-Loop for Adaptive Control

In some advanced cases, an AI might provide real-time control recommendations for a remediation system (like adjusting extraction rates if a plume is approaching a boundary). We require a human-in-the-loop in such decisions:

- The AI might recommend adjusting pump rates or injection volumes after analyzing data and model results. The operator/engineer reviews that suggestion, confirms it will not cause containment loss or other side effects, then implements it gradually.
- A safe approach is to pre-define safe bounds (never drop below a minimum flow or above a maximum level) and let AI tune within those bounds, with fail-safes that override to ramp pumps up if the plume moves unexpectedly.

The system might operate like adaptive cruise control — the engineer sets the parameters and constraints, monitors actively, and retains authority to override at any moment.

4.4 Ensuring Safe and Compliant Outcomes

No matter how sophisticated the predictive model or control suggestion, final remedial decisions are anchored to compliance metrics: meeting cleanup levels by a deadline, preventing offsite migration. AI can show an “optimal” path to reach goals, but if that path risks momentarily exceeding a local standard or fails to address a small ecological hotspot, it will not be acceptable. The engineer ensures any recommended strategy meets all regulatory criteria, often incorporating safety factors to account for model uncertainty, and communicates any AI-derived plan clearly in the remedial design documentation, including its limitations and how it was verified.

Module 5: Risk Prioritization & Decision-Making Metrics

Overview: At complex contaminated sites — or portfolios of many sites — deciding how to allocate resources often involves risk prioritization. AI can refine the metrics and analysis of risk by better predicting probability or severity of issues. Final risk decisions involve value judgments and regulatory frameworks that only human professionals and authorities handle; AI’s role is to provide evidence, not to set priorities.

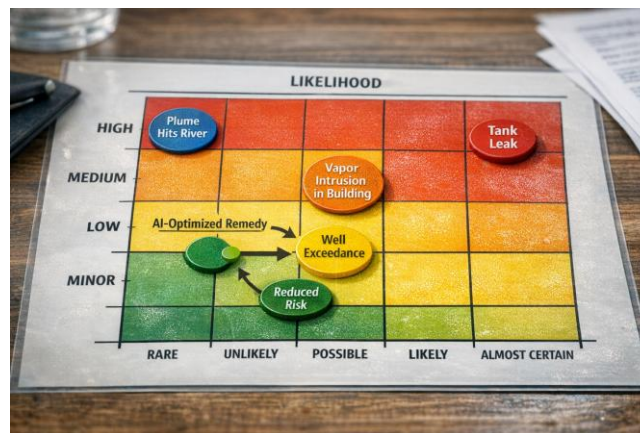
5.1 Risk Metrics & ML Inputs

In risk assessment, common metrics include:

- **Exceedance probability:** The chance a contaminant exceeds a threshold at a receptor. ML can improve these estimates by analyzing a wide range of conditions through Monte Carlo simulation. The engineer may present a risk as “<5% probability of exceeding the drinking water standard in the next 5 years, if current trends continue.”
- **Risk index or score:** Combining hazard, exposure, and remediation difficulty. AI can help quantify some components (like expected time to remediate based on modeling, or predicted future exposure if left unaddressed). AI does not decide the weighting of factors — regulators or organizational policy do — but AI ensures the underlying data feeding the index is as accurate and current as possible.

5.2 Prioritizing Actions or Sites

AI can assist by ranking scenarios based on predicted outcomes: for example, an algorithm might model the consequence of not remediating each site, finding that one site’s plume could reach a community well within ten years (serious) while another site’s plume might naturally attenuate before reaching receptors. The engineer confirms these models with conservative assumptions and ensures other factors (like cost and stakeholder concerns) are considered in the final priority decision. In adaptive management at a single site, ML can signal diminishing returns from a current method and suggest pivoting strategies; the decision to pivot is made in a risk-based review by engineers and regulators, using the ML evidence as part of the justification.



5.3 Handling Uncertainty in Risk Decisions

AI models inherently have uncertainty. When using them in risk contexts, engineers incorporate conservative biases:

- Use a high-end quantile of the ML prediction distribution (e.g., the 95th percentile of predicted concentrations) when calculating risk, to avoid underestimating exposure.
- In multi-criteria decisions, treat AI output as one input among many. Even if ML indicates a site's risk is low, regulators might identify an environmental justice concern making it a higher priority.
- Present a range of scenario outcomes rather than a single number: e.g., “depending on biodegradation rate, risk spans moderate to high — thus we plan as if the worst-case scenario to be safe.” AI helps generate these scenarios quickly, but the decision to lean toward worst-case is a conscious, conservative choice by the PE to guard against optimism bias.

5.4 Documentation & Regulatory Acceptance

If AI was used to support risk prioritization, it must be clearly documented:

- Detail assumptions (e.g., “model assumed no new sources and continuous monitoring of infiltration”).
- Note model accuracy and validation so regulators know how much trust to place in the outputs.
- Explain decisions: “We prioritized area X because ML projections showed potential off-site impact under conservative assumptions, whereas area Y's plume remains contained with minimal exposure risk.” This ensures that prioritization is seen as data-informed and hinged on rational, defensible criteria — not a black-box algorithm decision.

Module 6: Governance, Model Validation & Cybersecurity

Overview: Trusting AI in environmental projects demands robust governance. This module ties together how to ensure each AI model used is validated, controlled, and secure, and that its outputs are reliably integrated into formal project processes. We discuss model validation, change management, and cybersecurity aspects — since data integrity is key: an attacker altering sensor data could mislead an AI system if not properly protected.

6.1 Model Verification & Validation (V&V)

Before deployment, each AI model must be verified (does it compute as intended without errors) and validated (does it accurately represent reality within required tolerances). For example:

- **Benchmarking:** Run the model on scenarios where analytical or field results exist and check that outputs match within acceptable error, often set by policy (e.g., $\pm 10\%$ — and where bias exists, it should lean conservative rather than optimistic).
- **Cross-validation:** Use withheld data or independent test sets to ensure the model generalizes (common ML practice and often required by QA plans).
- **Peer review:** A second experienced engineer reviews the model methodology and results, akin to checking a complicated spreadsheet, to sign off that it is fit for its intended use.

6.2 Configuration Control & Change Management

We treat any AI model as a configured item:

- **Version control:** Each model version is uniquely labeled and stored. No modification to a model's code or training occurs without documentation. For instance, if we retrain a model with new monitoring data the following year, that becomes model v2.1, with release notes (what changed, how tested).
- **Approval for changes:** If model outputs feed a compliance decision, any significant model change should be reviewed and approved by the project's senior engineer or modeling lead as part of a Management of Change (MOC) procedure.
- **Traceability:** Maintain data lineage — what data was used to train each model version and where it came from. If environmental regulators ask “Did you include the last sampling round in your analysis?”, documentation must allow that question to be answered precisely.



6.3 Data Security & Integrity

Because AI relies on accurate data, cybersecurity is essential:

- **Sensor data authenticity:** If using IoT or remote sensors feeding an AI anomaly detector, protect those data feeds (encryption, authentication) so nobody can spoof a false safe reading that might trick the AI. A robust system may include cryptographic checks on field data so the AI rejects anything not properly authenticated.
- **Model security:** Restrict access to models and datasets. Only authorized team members should be able to modify the model or input dataset, requiring credentials and multi-factor authentication for remote access.
- **Backups & fail-safes:** Maintain backups of model and data (so a ransomware attack or server failure does not wipe monitoring or analysis capability). If an AI system goes offline, operations should revert automatically to manual monitoring and conservative assumptions by default.

6.4 Professional Defensibility & Communication

The engineer must be ready to defend any AI-assisted decision just as they would any traditional calculation. That means:

- **Keep audit trails:** If a regulator audits why a certain well was not sampled in a given quarter, and the rationale was that an ML analysis indicated low risk, the project team should produce the ML output, the engineer's review notes, and confirmation that sampling occurred in the subsequent round.
- **Communicate assumptions clearly:** In reports or public meetings, if AI methods were used, articulate in plain language what was done and the confidence level. For example: "We used a machine learning model, validated to 90% accuracy, to simulate many rainfall scenarios. It indicated a less than 1% chance of contaminant reaching the river. Out of caution, we still plan a contingency barrier installation if concentrations rise above a defined threshold — no decisions were automated without human oversight."
- **Maintain trust:** By documenting AI usage thoroughly, regulators and stakeholders see AI not as a mysterious black box making life-and-death calls, but as a sophisticated tool that the engineering team controlled and documented to improve environmental outcomes.

References

1. Liu, Y., et al. (2024). Machine Learning-Based Surrogate Modeling for Groundwater Contaminant Transport and Remediation Optimization. *Water*, 16(19), 2861. <https://www.mdpi.com/2073-4441/16/19/2861>
2. Pacific Northwest National Laboratory (PNNL). (2023). *Integrating Analytical Solutions and U-Net Model for Predicting Groundwater Contaminant Plume Changes Under Pump-and-Treat Operations*. U.S. Department of Energy. <https://www.pnnl.gov/publications/integrating-analytical-solutions-and-u-net-model-predicting-groundwater-contaminant>
3. National Society of Professional Engineers (NSPE). *Code of Ethics for Engineers*. <https://www.nspe.org/resources/ethics/code-ethics>
4. International Organization for Standardization. *ISO 9001:2015 — Quality Management Systems: Requirements*. Geneva: ISO, 2015.