

Engineering Ethics in the Age of Artificial Intelligence

Author: Dr Jeelan Moghraby, PhD., MBA

PDH Credits: 3

Course Overview

Course Description

This course examines how professional engineers should apply the ethical canons established by professional organizations, such as the National Society of Professional Engineers (NSPE), when using AI-assisted technologies in professional practice.

Artificial intelligence (AI) is increasingly integrated into engineering workflows, influencing analysis, modeling, drafting, design generation, simulation, and decision support. These tools extend analytical capability in ways that were impractical even a decade ago. However, the use of AI does not replace the professional judgment, legal accountability, or ethical obligations of the licensed engineer.

This course addresses the practical responsibilities that arise when AI systems influence engineering analysis and design: verifying automated outputs, maintaining human oversight, documenting AI-assisted workflows, and managing professional risk in environments where computational tools operate faster, and at greater scale than any individual engineer can directly supervise. Case studies drawn from documented engineering practice and adjacent technical fields illustrate how the profession's core ethical principles apply when automated systems are part of the analytical process.

Learning Objectives

Upon successful completion of this course, participants will be able to:

- Apply established engineering ethical canons when using AI tools in professional practice.
- Recognize the continuing professional responsibility of licensed engineers for all AI-assisted engineering outputs.
- Identify ethical risks associated with automated design and AI-generated analytical results.
- Apply verification and validation procedures for AI-assisted engineering analysis.
- Document AI-assisted design methods in a transparent and professionally accountable manner.
- Analyze the liability and risk management implications of AI use in engineering workflows.
- Integrate professional ethical frameworks into AI-supported engineering practice at the individual and organizational level.

Intended Audience

This course is intended for licensed professional engineers across disciplines including civil, structural, environmental, electrical, mechanical, aerospace, geotechnical, industrial, and systems engineering, as well as engineering professionals who use or anticipate using AI technologies within engineering workflows. The course is particularly relevant for engineers

involved in computational modeling, simulation, design automation, digital engineering environments, and AI-assisted analytical systems.

Contents

Professional Responsibility & Ethical Canons	6
1.1 Professional Responsibility and Ethical Canons in AI-Assisted Engineering	6
1.2 The Engineer as the Responsible Professional of Record	7
1.3 AI as a Professional Tool, and not a Professional Substitute	7
1.4 Applying the Public Safety Obligation When AI Influences Decisions	8
1.5 Honesty, Transparency, and the Prohibition on Misrepresentation	8
1.6 Professional Competence and Responsible Use of Advanced Technology	9
1.7 How Established Ethical Frameworks Apply to Emerging Technologies	9
Verification of AI-Generated Results	11
2.1 What Verification Requires in Practice	11
2.2 Independent Review and Peer Checking	12
2.3 Knowing When Not to Use AI.....	12
2.4 Verification Across Engineering Disciplines	13
Engineering Judgment in Automated Environments.....	15
3.1 What Independent Judgment Looks Like	15
3.2 Maintaining Human Oversight of Models and Simulations	16
3.3 Ethical Responsibility When Approving AI-Assisted Engineering Documents.....	17
3.4 Resisting Organizational Pressure to Delegate Judgment.....	17
Transparency, Documentation, & Professional Accountability.....	19
4.1 What Documentation of AI-Assisted Work Should Include.....	19
4.2 Traceability in AI-Assisted Workflows	20
4.3 Ethical Reporting of Analytical Limitations.....	20
4.4 Practical Documentation Standards for AI-Assisted Projects	21
Managing Professional Risk.....	22
5.1 Liability and the Limits of Technological Delegation	22
5.2 Data Quality and Intellectual Property Risks.....	23
5.3 Organizational Risk Management.....	23
5.4 The Emerging Regulatory Landscape	24
Ethical AI Practices in Engineering Organizations.....	25
6.1 Leadership in Technology Adoption.....	25
6.2 Internal Controls and Training.....	25
1.3 Maintaining Professional Standards as Engineering Technologies Evolve	26

6.4 Building AI Literacy Across the Engineering Team.....	26
Learning form AI Failures Across Industries	28
Case Study 1: Automation Bias and the Boeing 737 MAX.....	28
Case Study 2: IBM Watson for Oncology and Narrow Training Data.....	29
Case Study 3: Samsung Electronics and Profession Confidentiality.....	29

Chapter 1

Professional Responsibility & Ethical Canons



1.1 Professional Responsibility and Ethical Canons in AI-Assisted Engineering

Engineers have always relied on powerful tools to assist their work, from the slide rule to finite element analysis software. Each generation of technology raises the key question of who is responsible when the tool does a lot of the thinking. Artificial intelligence (AI) pushes this question further than ever before. AI systems can generate structural designs, predict geotechnical behavior, optimize electrical distribution networks, and draft regulatory compliance analyses faster than any individual engineer. The efficiency gains are impressive, but the ethical risks are significant as well.

What remains unchanged, no matter how advanced the tools become, is where professional responsibility lies. The licensed professional engineer is still the person accountable for the safety, integrity, and reliability of engineered projects. This principle reflects the social contract that the engineering licensure represents; society grants the engineer the authority to certify designs because the engineer accepts a personal and professional duty to protect the public.

This agreement with society is clearly stated in the codes of major engineering professional organizations, including:

- The National Society of Professional Engineers (NSPE)
- The American Society of Civil Engineers (ASCE)
- The American Society of Mechanical Engineers (ASME)
- The Institute of Electrical and Electronics Engineers (IEEE)

While the specific wording may differ for each organization, the order of priorities is consistent: public welfare comes before client interests, and professional judgment takes precedence over computational convenience.

1.2 The Engineer as the Responsible Professional of Record

AI systems can generate engineering calculations, perform simulations, or suggest design alternatives. It is important for the engineer to recognize that these outputs represent computational assistance rather than authoritative engineering conclusions, and their use does not reduce or transfer professional responsibility. In some respects, the presence of highly automated tools increases the importance of careful engineering oversight, because the gap between what a system produces and what a human expert would independently conclude can be harder to detect.

Licensed engineers should approach AI tools as they would any other engineering instrument or analytical method. Just as an engineer must understand the limitations of structural modeling software or hydraulic simulation programs, the engineer must understand the limitations and assumptions inherent within AI systems. Ethical engineering practice requires that the professional engineer remains as the ultimate authority in interpreting results and making final design decisions.

Professional engineering codes have long required that engineers perform services only within their areas of competence (NSPE, 2019; ASCE, 2017). Applying this requirement to AI will, therefore, require that the engineer develops sufficient familiarity with the tools they use to evaluate whether outputs are credible. An engineer who approves a design generated by an AI optimization platform without understanding the objective function, the constraints, or the training data, is not exercising competent professional judgment.

1.3 AI as a Professional Tool, and not a Professional Substitute

Understanding what AI can and cannot do is the starting point for applying engineering ethics. AI systems, whether based on machine learning, neural networks, or generative algorithms, are fundamentally pattern-recognition and optimization systems. They derive outputs from training data and objective functions defined by their developers. For example, an AI system trained on historical structural designs may suggest a configuration that optimizes material usage against some defined cost function, but it has no awareness of site-specific constraints, local code amendments, or the practical realities of fabrication in a given region. It cannot weigh safety margins against regulatory minimums with the contextual understanding a licensed engineer brings to that decision.

When finite element analysis software became widely available in structural engineering practice during the 1980s and 1990s, engineers remained responsible for defining boundary conditions, selecting realistic material properties, interpreting output with reference to physical behavior, and recognizing when a model was producing results that made no physical sense. The same logic applies to AI-assisted tools. The engineer must understand what the tool is doing, why it produces the outputs it does, and where its limitations lie.

Engineering design requires a combination of analytical reasoning, practical experience, regulatory knowledge, and ethical responsibility. AI systems cannot exercise professional judgment, cannot assume legal responsibility for design outcomes, and they cannot evaluate

complex trade-offs involving safety, environmental impact, economic constraints, and long-term system performance. These evaluations rest with the engineer.

1.4 Applying the Public Safety Obligation When AI Influences Decisions

When AI systems influence engineering decisions, the engineer must evaluate the numerical outputs as well as whether the analytical framework that produced those outputs is appropriate for the problem at hand.

AI optimization tools frequently prioritize efficiency metrics, such as cost reduction, weight minimization, or energy performance that often, but not always, align with good engineering outcomes. For example, an AI-driven structural optimization platform may propose reducing member sizes to achieve target material quantities. The proposed sizes may satisfy the governing equations in the training dataset while, under certain loading scenarios not well-represented in that dataset, they produce safety factors below what sound engineering practice requires. No analytical tool, however sophisticated, can substitute for the engineer's responsibility to evaluate whether the output is physically realistic and safely conservative under real-world conditions.

System-level thinking presents another dimension of the public safety obligation that AI tools often do not capture well. Infrastructure systems operate over decades. They are subject to deferred maintenance, changing operational demands, unusual loading events, and interactions with adjacent systems that were not anticipated at the time of design. An AI recommendation optimized for initial performance may not account for long-term durability or the consequences of component failure. Evaluating AI-generated recommendations within the context of long-term system reliability, as well as public welfare, is part of the engineer's responsibility.

1.5 Honesty, Transparency, and the Prohibition on Misrepresentation

Engineering codes are consistent in their requirement for honest communication of technical information. Engineers must not deceive clients, regulators, or the public about the nature and reliability of engineering analysis. However, AI-assisted workflows create specific risks to this.

When an AI system generates a design alternative or an analytical output, that output can appear highly authoritative. Complex visualizations, extensive numerical precision, and sophisticated-looking reports can create an impression of rigor that discourages scrutiny. An engineer who presents AI-generated analysis to a client without clearly communicating its basis, its limitations, and the degree to which it has been independently reviewed, is creating a misleading impression, even if no intentionally false statement has been made.

Transparency requires that engineers understand and can explain the analytical process behind their conclusions. If an AI system performs predictive modeling or generative design, the engineer should be able to describe the general methodology, identify the primary data inputs, and articulate the key assumptions embedded in the model. The NSPE Code of Ethics is explicit that engineers shall be objective and truthful in professional reports, statements, or testimony (NSPE, 2019). Presenting AI-generated analysis as though it were their own, without disclosing the role of the automated system, may not satisfy this requirement.

Misrepresentation can also arise from omission. If an engineering decision was significantly influenced by AI-generated modeling or predictive analysis, professional ethics requires that this influence be acknowledged in the engineering record and, where relevant, in communications with clients and regulators. Professional integrity requires that engineers maintain control over the engineering decision-making process, and use AI as a support tool that enhances analytical capability rather than one that determines final conclusions.

1.6 Professional Competence and Responsible Use of Advanced Technology

AI systems designed for engineering applications are often developed by software engineers or data scientists who may not have expertise in civil, mechanical, electrical, or environmental engineering. As a result, these tools may incorporate assumptions that do not fully align with real-world engineering conditions. Therefore, engineers will need to evaluate whether the analytical framework of a given AI system is compatible with the specific engineering problem being addressed, not just whether the tool is technically capable of producing a result.

Responsible use of AI requires engineers to understand the general capabilities and limitations of the technologies they employ. This includes familiarity with the input requirements, data dependencies, and computational scope of the AI system, at least sufficient enough to determine whether the resulting outputs are credible. The ethical requirement of competence also extends to continuing education. As AI technologies evolve, engineers must remain informed about emerging tools and their implications for professional practice. Using advanced computational systems without adequately understanding them, may introduce professional risks that the engineer has not adequately assessed. The engineer who uses AI tools well is one who maintains accountability, transparency, and competence in applying these tools as extensions of professional capability, rather than substitutes for it.

1.7 How Established Ethical Frameworks Apply to Emerging Technologies

A reasonable question arises when engineers first consider applying existing codes of ethics to AI-assisted practice: were those codes written with this technology in mind? The answer is no. The NSPE Code of Ethics dates in its present structure to 1964, with subsequent revisions. The ASCE and ASME codes similarly predate modern machine learning by decades.

Engineering codes of ethics are principles-based rather than technology-specific. They establish obligations related to competence, honesty, public safety, and professional accountability that are intended to govern engineering practice across its full scope of evolution. The NSPE Code does not specify how a licensed engineer must approach finite element analysis, geotechnical modeling software, or computational fluid dynamics platforms either. The expectation, consistently upheld by licensing boards and professional organizations, is that the ethical principles apply to whatever tools and methods engineers use, including AI.

The American Bar Association, facing analogous questions about AI-assisted legal work, concluded in a 2012 formal opinion that the duty of competence requires lawyers to

understand the benefits and risks of relevant technology. The National Academy of Engineering has similarly observed that professional responsibility frameworks are designed to adapt to technological change, not to be replaced by it (Russell & Norvig, 2020). Engineering ethics boards have begun reaching parallel conclusions: the engineer's obligation to exercise judgment, verify work, and protect the public remains paramount.

It is worth noting that engineering ethical frameworks are not static. Professional organizations including NSPE, ASCE, ASME, and IEEE have each started to examine how their existing guidance applies to AI-assisted practice, and several are actively developing supplementary guidance.

Chapter 2

Verification of AI-Generated Results



Engineers have always been expected to check their work, confirm that outputs from computational tools are consistent with physical reality, and exercise professional judgment in interpreting analytical results. The use of AI generated results makes this obligation more important and, in some respects, more difficult to satisfy.

It is more important because the scale of AI output is large. A generative design algorithm might produce dozens of structurally distinct alternatives, and each alternative requires evaluation. The pressure to accept results uncritically, simply because they arrived quickly and appear well-supported, runs directly against the engineer's verification obligation.

And it is more difficult because AI systems, particularly those based on machine learning, frequently operate as what practitioners call black boxes. The system receives inputs and produces outputs through internal processes that are not directly visible to the user. Unlike a finite element model, where the engineer can inspect mesh density, boundary conditions, and load cases, a neural network's internal reasoning is not available for direct review. This creates a genuine challenge for the traditional engineering approach of tracing a result back to its analytical foundation.

2.1 What Verification Requires in Practice

The opacity of AI systems changes how the verification requirement is done. Engineers using AI-assisted analytical tools should apply a layered approach to verification, working through the results at multiple levels of scrutiny.

The first level is an **order-of-magnitude check**. Does the result make physical sense? An AI-generated beam load, a pipe flow rate, or a thermal resistance value should be checked against simplified hand calculations or established rules before it is accepted. If a machine learning model predicts that a column will carry ten times the load indicated by a standard steel section capacity calculation, that discrepancy must be investigated before the design proceeds. The burden of explanation falls on the engineer; it should not be assumed that the AI result is correct.

The second level is **sensitivity analysis**. How does the output change when inputs vary within realistic ranges? AI models trained on specific datasets may perform reliably within those ranges but produce unrealistic results outside them. By varying key input parameters and observing whether the model's behavior remains physically plausible across the range of uncertainty, will enable engineers to develop a more informed view of where the tool's reliability begins to diminish.

The third level is **cross-verification** against traditional methods. For many engineering applications, it is practical to conduct parallel analysis using conventional software or classical analytical approaches. Where AI-generated results and traditional methods produce consistent conclusions, confidence in the AI output is strengthened. Where significant divergence exists, the cause must be identified before either result can be accepted.

2.2 Independent Review and Peer Checking

Many engineering organizations, and many applicable standards, require independent checking of calculations for safety-critical work. The presence of AI tools in the workflow does reinforce this. A reviewer who understands that the primary analysis was AI-assisted should approach the check with awareness of the specific failure modes of machine learning systems: sensitivity to training data quality, extrapolation behavior, and the risk that statistical correlations in the model do not reflect causal physical relationships.

The National Institute of Standards and Technology (NIST) Artificial Intelligence Risk Management Framework identifies several dimensions of trustworthiness that are directly relevant to engineering practice: reliability, safety, explainability, and the ability to validate model performance against known benchmarks (NIST, 2023). Engineers and engineering organizations can draw on this framework when designing review procedures for AI-assisted work. The framework does not prescribe engineering-specific protocols, but its categories of risk align well with the verification challenges that engineers face.

Independent checking is particularly important when AI tools influence the design of systems that involves public safety. For example, a machine learning model used to inform the design of a stormwater detention basin, a structural connection detail, or an electrical protection relay scheme, is influencing outcomes that could affect people if they fail. The standards of review applied to those outputs should reflect this consequence, and not merely the analytical method that was used to generate them.

2.3 Knowing When Not to Use AI

It is important for engineers to understand whether an AI tool is appropriate or not for a given problem. Using an AI tool because it is available, rather than because it is appropriate,

is not good professional judgement. Selecting analytical methods is itself an engineering decision, and it carries the same professional accountability as any other decision made during a project.

Some engineering challenges may involve rare events, novel conditions, or regulatory requirements unlikely to be well-represented in any historical training dataset. In such cases, traditional engineering analysis may offer more defensible conclusions precisely because its assumptions are explicit and can be directly evaluated.

2.4 Verification Across Engineering Disciplines

While the principles of verification are consistent across engineering practice, their application differs by discipline. The following discipline-specific examples illustrate how this may occur in practice:

- *Structural engineering:* AI-assisted design tools may generate member sizes, connection details, or framing configurations. Verification of the results should confirm that governing load combinations have been applied, that applicable code provisions have been satisfied, and that the design reflects the correct geologic, wind, and seismic hazard parameters for the project location. An AI tool trained on projects in one region may not incorporate the seismic detailing requirements of another, and this is not necessarily visible in the output.
- *Geotechnical engineering:* AI tools are increasingly used to predict bearing capacity, settlement, and slope stability from site investigation data. These predictions are particularly sensitive to data quality and training domain. Verification should assess whether the soil conditions at the project site are representative of the conditions the model was trained on, and whether the model produces physically plausible results when compared against classical analytical methods, such as Terzaghi bearing capacity equations or Bishop's method for slope stability. Significant divergence warrants investigation before AI-generated geotechnical recommendations are adopted.
- *Environmental engineering:* This discipline has its own verification challenges, particularly for AI tools used to predict contaminant transport and fate. Contaminant behavior is highly site-specific and sensitive to parameters including soil heterogeneity, groundwater chemistry, and microbial activity, that may not be well-represented in historical training datasets. Engineers verifying AI-assisted contaminant modeling should assess whether the site conditions align with the model's training domain. They should also apply simplified analytical screening calculations to confirm that predicted contaminant concentrations and migration rates fall within physically plausible ranges.
- *Electrical engineering:* AI tools may assist with load forecasting, system optimization, or fault detection. Verification should confirm that protection coordination, load flow, and fault current calculations are consistent with established power systems analysis. It should also confirm that the AI-generated system configurations satisfy applicable standards including IEEE and the National Electrical Code requirements. Automated recommendations that alter system protection settings or equipment ratings require

particularly careful independent review, given the safety consequences of protective relay miscoordination.

Across all disciplines, the underlying verification obligation is the same: the engineer must be able to demonstrate that AI-generated outputs were evaluated against the physical and regulatory requirements applicable to the specific project, and where not simply accepted because they were produced by a credible-looking system.

Chapter 3

Engineering Judgment in Automated Environments



Automation bias is well-documented. It describes the tendency to favor outputs from automated systems over independent judgment, even when the automated output is incorrect and the person has information that would reveal the error (Parasuraman & Riley, 1997). It has been implicated in aviation accidents, medical misdiagnoses, and financial system failures, and is now a recognized risk factor in engineering.

The mechanism of automation bias is simple. When an AI system produces a result confidently, with apparent mathematical precision and extensive supporting output, questioning it requires effort. It means conducting additional calculation, consulting additional sources, and potentially disrupting a workflow that has moved quickly. Engineers working under time pressure or project budget constraints may find it easy to rationalize acceptance because the result looks right and the system is sophisticated. These rationalizations are precisely the form of judgment failure that professional ethics guards against.

3.1 What Independent Judgment Looks Like

Independent judgment requires approaching AI outputs with the same analytical mindset that good engineers apply to any source of technical information: critical, informed, and aware of the source's limitations.

In practice, this involves asking a structured set of questions about each significant AI-generated result:

- What data trained this model, and is that data representative of this project's conditions?
- What objective function does the model optimize, and does that objective fully capture the engineering requirements of this application?
- Are there safety, durability, or regulatory considerations that the model may not have adequately weighted?
- Has the result been cross-checked by any independent method?

For many routine applications where AI tools are well-established and the engineer is familiar with their behavior, these questions can be addressed efficiently without adding significantly to project timelines. However, they must be addressed, and the answers must reflect genuine professional engagement rather than a reflexive endorsement of automated output.

3.2 Maintaining Human Oversight of Models and Simulations

Human oversight in AI-assisted engineering means maintaining awareness of how automated systems are being used throughout the analytical process, understanding what assumptions are embedded in the models, and ensuring that the engineering intent of the analysis is reflected in the computational setup.

An AI system given poorly specified inputs will produce poorly founded outputs. This is easy to overlook in practice when the tool's interface makes it straightforward to run an analysis without clearly prompting the user to define and defend the problem setup. Engineers must take responsibility for the quality of the inputs they provide to AI systems, for the appropriateness of the problem framing, and for the interpretation of results in the context of the actual engineering problem being solved.

This obligation is especially significant for AI tools that interface with Building Information Modeling platforms, digital twin systems, or automated code-checking software. These environments may involve chains of automated analysis where the output of one system becomes the input to the next, and where assumptions introduced at the beginning of the chain propagate through every subsequent step. Engineers overseeing such workflows must understand where critical assumptions enter the process and verify that those assumptions are technically defensible.

AI models may occasionally produce results that appear mathematically valid but reflect unrealistic or impractical engineering solutions. Optimization algorithms may propose design configurations that minimize cost or maximize performance metrics, while ignoring practical considerations such as constructability, maintenance requirements, or long-term durability. Human oversight ensures that these broader engineering considerations remain part of the decision-making process.

3.3 Ethical Responsibility When Approving AI-Assisted Engineering Documents

When an engineer signs and seals an engineering document, that act has legal and professional consequences. The seal represents the engineer's personal certification that the work meets professional standards, and that the engineer accepts responsibility for its accuracy and integrity. This responsibility does not change when the analysis underlying the document was produced in whole or in part by an AI system.

Before approving AI-assisted engineering documents, engineers should be able to confirm four things:

1. They have reviewed the key analytical results.
2. They understand how those results were produced.
3. Independent verification has been conducted where required by applicable standards or organizational protocols.
4. The conclusions are consistent with engineering principles and applicable codes.

An engineer who cannot answer these questions for a document they are being asked to seal has not completed the professional review that approval represents. There is no comfortable middle ground. Either the engineer has done the work to satisfy themselves that the design is sound, or they have not. AI-assisted analysis may make certain aspects of design development faster and more capable, but it does not change what professional approval means or what it requires.

3.4 Resisting Organizational Pressure to Delegate Judgment

Engineering projects rarely fail because a single engineer made a bad individual decision in isolation. It is more common that they fail because organizational conditions created incentives that made sound professional judgment difficult to maintain. Schedule pressure, fee compression, client demands for rapid turnaround, and internal cultures that treat technical review as a cost center rather than a professional requirement, are well-documented contributors to engineering failures long before AI system were developed.

AI-assisted workflows can amplify these pressures in specific ways. Because AI tools can generate results quickly, there is a temptation to interpret speed as thoroughness. A platform that produces a structural design in minutes may appear to have done the analytical work that justifies that result, even when the verification that would justify it has not been completed. Engineers in organizations where AI tools have been adopted often report that the implicit expectation shifts from 'has this been properly reviewed?' to 'the platform produced it, so it should be fine.' That shift is a professional risk.

Engineers who raise concerns about inadequate review of AI-generated outputs in their organizations may encounter resistance, particularly if the organization has made a significant investment in adopting a particular tool. Professional ethics codes address this directly. The NSPE Code of Ethics states that engineers shall not be influenced in their professional duties by conflicting business or personal interests, and that engineers shall act in such a manner as to uphold and enhance the honor, integrity, and dignity of engineering

(NSPE, 2019). Raising a legitimate concern about insufficient verification of AI-generated design outputs is an act consistent with these obligations.

Senior engineers and engineering managers have a particular responsibility here. Organizational culture around technical review is set from the top. If leadership treats verification of AI-generated outputs as an obstacle to be minimized rather than a professional obligation to be met, engineers throughout the organization will receive that signal clearly. Building a culture in which adequate review of AI outputs is expected and supported is a leadership function that cannot be delegated to a policy document.

Chapter 4

Transparency, Documentation, & Professional Accountability



Engineering documentation exists to create a record of what was decided, why it was decided, who made the decision, and on what basis. That record serves the project in the short term and the profession in the long term. When something goes wrong, it is the record that allows engineers, investigators, and regulators to understand what happened. When something goes right, it is the record that allows lessons from the project to be transferred to future work.

AI-assisted engineering creates new documentation obligations. The core challenge is that AI systems, particularly machine learning systems, may not generate the kind of explicit, step-by-step analytical record that traditional engineering calculations produce. The engineer must create this record separately, drawing on their understanding of what the AI system did and translating it into the documentary form that professional accountability requires.

4.1 What Documentation of AI-Assisted Work Should Include

At a minimum, documentation of AI-assisted engineering work should identify the tool used, the version or release, the nature of the inputs provided, the types of outputs generated, and the review process applied to those outputs before they were incorporated into engineering decisions. This information allows a future engineer, reviewer, or investigator to evaluate the basis of the design.

Documentation should also record the key assumptions underlying the AI analysis. AI systems embed assumptions in their training data, objective functions, and model architecture. When those assumptions are relevant to the safety or performance of the

design, they should be identified in the engineering record. For example, if the AI model was trained on soil conditions that differ from the project site, or on structural loading scenarios that do not match the actual use case, those discrepancies should be noted and their implications addressed before the design is finalized.

Where AI-generated results were rejected or modified following engineering review, the documentation should explain why. An engineering record that shows only the final accepted results does not reflect the analytical process that led to them. It also does not create the traceability that professional accountability requires. The documentation should include the reasoning behind significant engineering decisions as part of the record.

4.2 Traceability in AI-Assisted Workflows

Traceability, the ability to follow the analytical pathway from input data to final design conclusion, is a core requirement of defensible engineering practice. It is also the part of documentation that AI-assisted workflows make most difficult to maintain as a machine learning model uses input data to generate a prediction through internal processes that are not directly auditable. Engineers must compensate for this by creating supplementary records that establish the analytical context around the AI-generated results. This means documenting the engineering basis for accepting or modifying AI outputs, the independent checks performed, and the professional reasoning applied in their interpretation of the results.

In projects involving public infrastructure, industrial facilities, or other safety-critical systems, the expectation of traceability should be treated as a design requirement from the outset. Documentation protocols should be established at the beginning of AI-assisted projects, specifying what will be recorded, when, and by whom.

4.3 Ethical Reporting of Analytical Limitations

Professional honesty requires that engineers report not only what their analysis found, but also what could not be determined with confidence. AI-generated predictions carry uncertainty that should be communicated clearly to clients, regulators, and other stakeholders who will rely on engineering conclusions.

An AI model that predicts structural performance based on historical data from similar structures, may produce well-founded results in typical cases and poorly founded results in atypical ones. Engineers must evaluate which situation applies to the project at hand. When the AI analysis is operating in conditions that differ from its training domain, the uncertainty of its predictions is likely higher, and that higher uncertainty should be communicated in engineering reports.

Overstating the confidence of AI-generated analysis is a form of misrepresentation, even if it is unintentional. Engineers have an obligation to understand the uncertainty characteristics of the tools they use, and to report analytical conclusions in terms that accurately reflect that uncertainty. The NIST AI Risk Management Framework explicitly identifies the communication of uncertainty as a component of trustworthy AI deployment (NIST, 2023), and this principle aligns directly with the ethical requirement of honest technical communication that governs engineering practice.

4.4 Practical Documentation Standards for AI-Assisted Projects

Establishing documentation standards before a project begins is one of the most effective steps an engineering organization can take to maintain accountability in AI-assisted workflows. However, AI tools may not automatically generate the records that traditional software and hand calculation methods produce as a byproduct of the analytical process.

A practical documentation standard for AI-assisted engineering work might include the following elements recorded at each stage where an AI tool contributes to a design decision:

1. **A tool identification record:** the name, version, and developer of the AI tool; the date on which it was used; and the nature of the engineering problem to which it was applied.
2. **An input record:** the data provided to the AI system, including its source, its date, and any preprocessing applied before submission to the tool.
3. **An output record:** the results generated by the AI system, including key numerical values, the form in which results were delivered, and any confidence intervals or uncertainty ranges provided by the tool.
4. **A verification record:** the independent checks performed on the AI-generated output, the methods used, and the engineer's professional conclusion regarding the reliability of the output.
5. **A decision record:** the engineering decision made based on the AI-generated and independently verified analysis, and the professional judgment applied in reaching that decision.

This level of documentation is achievable without disproportionate effort if it is built into project workflows from the start. Organizations that establish standard templates for each of these record types, integrated into their existing project management and quality assurance systems, find that the additional documentation burden is modest compared to the professional protection and organizational learning it provides.

The documentation should be treated as a living record, updated as the project progresses and as AI-generated results are reviewed, modified, or rejected. An engineering record that documents only the final accepted design, without showing the analytical pathway that led to it, does not meet the standard that professional accountability requires. When AI-generated outputs are rejected following review, the record of that rejection and the reasons for it are as professionally important as the record of what was accepted.

Chapter 5

Managing Professional Risk



Engineering risk management involves recognizing where the boundaries of reliable knowledge lie, establishing practices that manage exposure within those boundaries, and building organizational structures that support sound professional judgment. AI introduces new territory across all three of these dimensions.

The technical risks of AI-assisted engineering include model uncertainty, training data limitations, extrapolation behavior, and the challenge of verifying black-box outputs. What receives less attention is professional liability. An engineer whose design is later found to be deficient cannot point to the AI system as the responsible party. The AI system has no professional license, no legal standing, and no accountability while the engineer has all three.

5.1 Liability and the Limits of Technological Delegation

Professional liability in engineering is determined by whether the engineer exercised reasonable care, competence, and judgment in providing engineering services. Introducing AI tools into the analytical process does not alter this standard; it relocates the question of where reasonable care must be demonstrated. An engineer can no longer satisfy the standard simply by documenting what calculations were performed, if those calculations were generated by a system whose reliability the engineer cannot demonstrate.

Courts and licensing boards have not yet developed extensive case law on AI-specific engineering liability. This is predicted to change as AI tools become more widely used, and as failures attributable to AI-assisted analysis become more visible. Engineers and engineering organizations would be unwise to wait for that body of precedent to develop before establishing professional practices that can withstand its scrutiny.

It remains with the engineer to demonstrate that AI tools were selected appropriately, used within their validated domain, and independently verified before their outputs were incorporated into engineering decisions. That demonstration must be supported by documentation. Engineers and organizations that build that practice now will be in a substantially better position than those who do not.

5.2 Data Quality and Intellectual Property Risks

AI systems depend on data. The quality of the data shapes the reliability of the model. However, the data introduces professional risks that go beyond reliability: it also raises questions of intellectual property and client confidentiality.

When an engineer submits project-specific data to a cloud-based AI platform, that data may include proprietary site investigation results, client-confidential project parameters, or commercially sensitive design specifications. The terms of service governing many commercial AI platforms do not guarantee that submitted data will not be used to train or refine the AI model, or that it will not be accessible to third parties under certain conditions. Engineers should review the data handling policies of AI tools they use professionally. They should also consider whether submitting specific categories of client data is consistent with their confidentiality obligations under applicable professional codes and any contractual arrangements with clients.

This is an area where professional guidance is actively developing. Some state professional engineering boards have begun issuing guidance on AI use, and organizations including NSPE and ASCE are developing position statements. Engineers should monitor these developments and ensure that their practice remains aligned with emerging professional standards (NCEES, 2023).

5.3 Organizational Risk Management

Individual engineers cannot manage all the risks associated with AI-assisted engineering practice on their own. Organizational structures and policies are necessary to ensure that AI tools are selected, deployed, and reviewed in a manner consistent with professional obligations.

Engineering organizations adopting AI tools should establish clear policies covering several areas including:

- Which AI systems are approved for use in engineering work.
- What categories of engineering problems they may be applied to.
- What verification procedures apply to AI-generated outputs.
- What documentation is required when AI tools contribute to engineering decisions.

These policies should be developed with input from engineers who will use the tools, not solely by technology or management personnel, and they should be reviewed periodically as the capabilities and limitations of AI systems evolve.

Quality management systems provide a useful framework for integrating AI oversight into organizational practice. Engineering organizations already operating under ISO 9001 or similar quality management frameworks can extend their existing quality control

procedures to cover AI-assisted workflows. This integration is more likely to produce consistently applied controls than standalone AI governance policies that exist separately from an organization's established technical review processes.

5.4 The Emerging Regulatory Landscape

The regulatory environment for AI in engineering is in its early stage of development. Engineers who treat this as a distant concern are likely to find themselves behind the curve when state licensing board guidance, federal agency requirements, or client contractual demands begin to specify how AI tools must be used and documented in engineering work.

At the federal level, the NIST AI Risk Management Framework, published in 2023, represents the most developed regulatory reference for AI governance in technical practice. While not legally binding, it is increasingly referenced by federal agencies in their AI adoption guidelines, and it provides a structured approach to evaluating AI risk that engineering organizations can adapt for their specific contexts (NIST, 2023). Federal procurement requirements for engineering services on federally funded projects may increasingly reference the AI Risk Management Framework or incorporate similar requirements as AI use in engineering becomes more prevalent.

The European Union AI Act, finalized in 2024, establishes a risk-based regulatory framework for AI systems that applies to EU member states, and has implications for engineering firms operating internationally. Under the Act, AI systems used in critical infrastructure, safety-critical product components, and infrastructure management are classified as high risk. Systems in this category are subject to requirements including transparency, human oversight, and technical documentation; requirements that overlap significantly with the verification and documentation obligations discussed in this course. Engineers at firms with international operations should monitor how the EU AI Act affects their practice.

At the state level, engineering licensing boards are beginning to address AI explicitly. Several boards have issued informal guidance or position statements noting that existing professional responsibility requirements apply fully to AI-assisted work. The National Council of Examiners for Engineering and Surveying (NCEES) has signaled that its model rules and competency standards apply to the use of AI tools in engineering practice, consistent with the broader principle that professional obligations are not tool-specific (NCEES, 2023). Engineers should check with their state licensing boards for current guidance, as this is an area where board positions are constantly evolving.

Professional organizations, including the NSPE, ASCE, ASME, and IEEE, have each initiated work to develop more specific guidance on AI in engineering practice. The ASCE's Committee on Technical Advancement and several ASME technical divisions have published working papers on AI governance in engineering. Staying current with that guidance, and contribution to the professional conversations that shape it, is part of responsible engineering practice in a period of rapid technological change.

Chapter 6

Ethical AI Practices in Engineering Organizations

Engineering leadership carries responsibilities that extend beyond technical direction. Engineers in senior roles, whether as principals of private practice, senior engineers within design firms, or technical leaders within public agencies, have an obligation to shape the professional culture of the organizations they lead. That obligation now includes how those organizations adopt and govern the use of AI.

This is primarily a question of ethics. The ethical frameworks that govern engineering practice apply to organizations as well as individuals. Senior engineers who allow AI tools to be used in ways that compromise professional judgment or place public safety at risk, are not fulfilling their leadership obligations.

6.1 Leadership in Technology Adoption

The adoption of AI tools within engineering practice can be driven by genuine capability improvements, by competitive pressure, or by enthusiasm for new technology that outpaces understanding of its limitations. The responsibility of leadership is to ensure that adoption is guided primarily by their improvement capabilities.

A useful discipline is to require that any AI tool proposed for use in engineering work be evaluated against a defined set of professional criteria before adoption:

- What is the tool's validated domain of application?
- What independent benchmarks support claims about its accuracy?
- What are its known limitations?
- What verification procedures are required to use it responsibly?

These questions do not hold AI tools to a higher standard than other engineering software. They apply the same standard that responsible engineering practice has always required.

Engineers in leadership positions must also set clear expectations about the role of AI tools within the organization. AI should function as an analytical aid that extends capability, not as a substitute for professional evaluation. When engineers observe that AI outputs are being accepted without adequate review, or that project timelines are being used to justify skipping verification steps, raising those concerns is a professional obligation.

6.2 Internal Controls and Training

Internal engineering controls for AI-assisted work should be integrated into existing quality management structures rather than treated as separate processes. This makes them more likely to be consistently applied and more easily auditable. Controls should cover tool approval, verification requirements, documentation standards, and review procedures, with appropriate calibration to the safety criticality of the application.

Training is equally important. Engineers should not be expected to supervise AI-assisted analytical processes without preparation. Training programs should cover how AI tools

operate at a general level, the failure modes most relevant to engineering applications, the verification and documentation obligations that apply, and the ethical responsibilities that remain with the engineer regardless of what the tool produces. This requires that engineers become professionally literate in the tools they are expected to use.

1.3 Maintaining Professional Standards as Engineering Technologies Evolve

Engineering has always been a profession shaped by technological change. Throughout the profession's history, engineers have adopted new analytical methods, materials, and computational tools to improve the performance and reliability of engineered systems. AI represents the next significant step in this evolution, and it is appropriate to approach it with informed engagement, and not reflexive resistance or uncritical acceptance.

The fundamental ethical obligations of the engineering profession remain constant through technological change. Engineers must continue to prioritize public safety, maintain honesty in professional communication, and perform engineering services with competence and diligence. These principles provide the foundation for responsible engineering practice regardless of the analytical tools involved. AI does not require a new set of ethics. It requires a disciplined application of the ethics the profession already has, to a context where the tools are more capable, the outputs are less transparent, and the temptation to delegate judgment is correspondingly greater.

The engineering profession has a history of adapting to major technological transitions. What distinguished the transitions that went well is the discipline the engineers brought to adopting it; understanding the new tools well enough to use them critically, holding to professional obligations when schedule and commercial pressures pushed in the opposite direction, and building organizations where quality was treated as a fixed requirement. This also applies to the use of AI technologies.

6.4 Building AI Literacy Across the Engineering Team

Technical literacy in AI is not uniformly distributed, and the gap between engineers who understand these tools well and those who do not creates specific risks within engineering organizations. An engineer who does not understand how a machine learning model derives its predictions cannot effectively review an output produced by one. Training that builds sufficient shared understanding across a team helps to maintain consistent professional standards when AI tools are in use.

Effective AI literacy training requires engineers to have a working understanding of several concepts that are directly relevant to professional practice. Engineers should understand:

- The difference between interpolation and extrapolation in AI models, and why models tend to be less reliable when applied to conditions outside their training domain.
- The concept of training data bias and how it can lead AI systems to produce recommendations that reflect the patterns of the past rather than the requirements of the specific project at hand.

- Why AI model outputs often lack the explicit analytical chain that traditional engineering calculations provide, and what that means for documentation and traceability.

Beyond conceptual training, engineers benefit from structured opportunities to apply verification skills to AI-generated outputs in low-stakes settings before they are expected to do so on live projects. Organizations can create internal exercises in which engineers are given AI-generated design outputs alongside the input data that produced them, and are asked to identify potential issues, conduct verification checks, and document their review. This kind of practice develops the habits that professional responsibility requires, and it surfaces tool-specific issues that can be addressed through protocol development before they appear on client projects. The NIST AI Risk Management Framework includes guidance on AI literacy that can be adapted for engineering training programs.

Chapter 7

Learning from AI Failures Across Industries

Various industries, including aviation, healthcare, and technology, have published documented cases of AI systems that were used without sufficient oversight, without adequate understanding of the AI's training data limitations, or without adequate consideration of the professional obligations that govern the use of sensitive information. These cases are instructive because they occurred under similar conditions found in engineering: high-stakes decisions, consequential outcomes, and professional accountability. In each case, technically sophisticated outputs created pressure to accept results without adequate scrutiny. The lessons learned from these cases, when applied to engineering, illustrate how the obligations of professional responsibility, verification, transparency, and judgment apply when AI or automated systems are part of the analytical process.

Case Study 1: Automation Bias and the Boeing 737 MAX

The case is not necessarily an AI case, but it illustrates what happens when engineers and the organizations they work for become insufficiently critical of automated system outputs, allow economic and schedule pressures to compress verification processes, and fail to maintain transparent documentation of system capabilities and limitations.

In the development of the Boeing 737 MAX, engineers and regulators approved a flight control system known as the Maneuvering Characteristics Augmentation System (MCAS). MCAS was an automated system designed to adjust the aircraft's pitch behavior under specific flight conditions. During its approval, the system's expanded authority over flight control inputs, and its reliance on data from a single angle-of-attack (AoA) sensor were not fully evaluated or communicated to the Federal Aviation Administration's aircraft evaluation group. In addition, this critical information about the system's behavior was not included in the pilot training materials (Defazio & Larsen, 2020). The flawed design of the AoA triggered repeated nose-down inputs that resulted in two crashes, five months apart, killing 346 people.

The main issue with the MCAS was that it had authority over the aircraft that exceeded what the pilots understood it to have. Critically, the design assumed that the pilots would respond in a way they had not been trained for, and the information necessary to correct that assumption was not passed on to the people who needed it.

This case has direct parallels in AI adoption. An AI tool that produces design outputs with apparent authority, where the internal logic is not transparent to the engineers reviewing it, and whose scope of influence on the final design has not been clearly defined, will result in a similar condition. The engineers are nominally in control, but the degree to which the automated system is actually shaping decisions may exceed what they understand it to be. The risk is not that the AI or automation is untrustworthy in absolute terms, but that the people responsible for it do not have a sufficiently clear picture of what it is doing and why.

Engineers reviewing AI-assisted designs should ask questions such as:

- What happens when the AI system receives anomalous inputs?
- What are the failure modes of the model, and how were they evaluated?
- Has the system been validated against conditions that differ meaningfully from its design environment?
- Are the engineers reviewing and approving the output fully informed of how the system works and what its limitations are?

Case Study 2: IBM Watson for Oncology and Narrow Training Data

IBM Watson for Oncology was an AI system designed to assist oncologist in identifying cancer treatment options. It did this by analyzing patient records and generating treatment recommendations with confidence ratings. The system was used in several hospitals globally and was promoted in its ability to process vast medical and clinical data at a scale beyond any individual physician's capabilities.

In 2018, STAT News reported that in multiple cases, the cancer treatment recommendations by IBM Watson were unsafe and incorrect (Ross & Swetlitz, 2018). It included suggestions that conflicted directly with established national clinical guidelines and was deemed inappropriate by the oncologists for their patients. The investigation revealed that IBM Watson was trained on a relatively small number of hypothetical patient cases, and the opinions of a few specialists. Because of this, it generalized its approach in ways that would not hold across different cancer types and clinical diversity of real patient populations.

This case shows that an AI system that has been trained on data that does not adequately represent the conditions to which it will be applied, can produce outputs that appear credible but are fundamentally wrong. Engineers evaluating AI tools for their use in design or analysis should always examine the data the model was trained on, and whether that data genuinely represents the engineering conditions they are applying it to.

Case Study 3: Samsung Electronics and Profession Confidentiality

In 2023, on three separate occasions, Samsung Electronics employees uploaded proprietary information to Open AI's ChatGPT to assist with work tasks, including debugging code and drafting documentation. The submissions included proprietary semiconductor source code, internal meeting notes, and hardware performance data. The engineers involved had not considered, or were not informed, that the data submitted to ChatGPT was saved on their servers, and could, under their terms of service at the time, be used to train future models. Samsung subsequently issued a ban on the use of generative AI tools by their employees (Ray, 2023).

This case illustrates a category of professional risk that many overlook. Engineers are bound by confidentiality obligations to their clients. Before using any cloud AI platform with client-specific project data, engineers should review the platform's data handling policies, assess whether those policies are consistent with their confidentiality obligations and any contractual restrictions, and, where necessary, seek explicit client consent for their use. When the engineer is not confident of these aspects, the appropriate response is to identify

alternative analytical approaches rather than to proceed and hope that the platform's practices are compatible with professional obligations.

Each of the cases discussed in this chapter share a common characteristic. The ethical failures they illustrate were not caused by the AI tools themselves. They were caused by insufficient professional engagement with what those tools were doing, what their limitations were, and what their professional obligations required in that context.

Engineering ethical frameworks are well-established, and should be applied with the same discipline in AI-assisted environments as in any other.

References

- American Society of Civil Engineers. (2017). Code of ethics. ASCE. <https://www.asce.org/career-growth/ethics/code-of-ethics>
- American Society of Mechanical Engineers. (2012). ASME code of ethics of engineers. ASME. <https://www.asme.org/about-asme/governance/ethics>
- European Commission High-Level Expert Group on Artificial Intelligence. (2019). Ethics guidelines for trustworthy AI. European Commission. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>
- Institute of Electrical and Electronics Engineers. (2020). IEEE code of ethics. IEEE. <https://www.ieee.org/about/corporate/governance/p7-8.html>
- National Council of Examiners for Engineering and Surveying. (2023). Model law engineer: Model rules. NCEES. <https://ncees.org/engineering/licensure/>
- National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
- National Society of Professional Engineers. (2019). NSPE code of ethics for engineers. NSPE. <https://www.nspe.org/resources/ethics/code-ethics>
- Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230–253. <https://doi.org/10.1518/001872097778543886> [Please verify before publishing]
- Ray, S. (2023, May 02). Samsung Bans ChatGPT Among Employees After Sensitive Code Leak. *Forbes*. <https://www.forbes.com/sites/siladityaray/2023/05/02/samsung-bans-chatgpt-and-other-chatbots-for-employees-after-sensitive-code-leak/>
- Ross, C., & Swetlitz, I. (2018, July 25). IBM's Watson supercomputer recommended 'unsafe and incorrect' cancer treatments, internal documents show. *STAT News*. <https://www.statnews.com/2018/07/25/ibm-watson-recommended-unsafe-incorrect-treatments/>
- Russell, S., & Norvig, P. (2020). *Artificial intelligence: A modern approach* (4th ed.). Pearson.
- Defazio, P.A., & Larsen, R. (2020). The design, development, and certification of the Boeing 737 MAX. U.S. House Committee on Transportation and Infrastructure. https://www.govinfo.gov/app/details/GOVPUB-Y4_T68_2-PURL-gpo144993

Disclaimer

This course was prepared by Dr Jeelan Moghraby, PhD., MBA.
Copyright © Dr Jeelan Moghraby 2026. All rights reserved.

Educational Use Only

The content of this course is provided for educational and professional development purposes only. While every effort has been made to ensure the accuracy and reliability of the information presented, MondoScience Courseware makes no representations or warranties regarding the completeness, accuracy, or applicability of the content to any specific professional context or jurisdiction. Users are advised to consult primary sources, applicable codes and standards, and exercise independent professional judgment when applying knowledge gained from this course.

PDH and Licensure

Completion of this course may qualify for Professional Development Hours (PDH) credit toward engineering license renewal, subject to the rules and requirements of the relevant state licensing board. It is the responsibility of the individual engineer to confirm eligibility and report hours in accordance with their jurisdiction's continuing education requirements.

References

References cited in this course have been selected to reflect current standards and published literature. Users are encouraged to verify the currency and accuracy of all references independently, particularly where standards or regulations may have been updated since publication.
