



AI-BASED SUPPLY CHAIN SECURITY FOR CONTAINERIZED APPLICATIONS

Author Name: Roshan Kakarla

Designation: DevOps Engineer

Department: Information Technology

Affiliation: Indiana Wesleyan University

ABSTRACT

The software supply chain underpinning containerized applications has emerged as one of the most critical and least understood sources of systemic risk in modern cloud-native systems. Container platforms enable high-velocity software delivery by composing applications from layered images, third-party dependencies, automated build pipelines, infrastructure-as-code definitions, and dynamic runtime orchestration. While this architectural model improves scalability and operational efficiency, it also creates expansive trust dependencies that extend far beyond the runtime perimeter. Existing supply chain security approaches rely predominantly on static verification, discrete scanning stages, or rule-based policy enforcement. These mechanisms provide limited visibility into how risk evolves over time and fail to capture the compounding effects of dependency drift, identity compromise, behavioral anomalies, and governance misalignment across the software lifecycle.

This paper introduces a novel, AI-based supply chain security framework that reframes supply chain protection as a continuous, risk-aware control problem rather than a compliance or tooling exercise. The core contribution is a governed supply chain security control plane that integrates probabilistic trust modeling, artifact lineage analysis, behavioral telemetry, and policy constraints to reason about supply chain risk across build, distribution, deployment, and runtime phases. Artificial intelligence is employed not as an autonomous decision-maker, but as an inference mechanism that synthesizes heterogeneous signals to identify trust decay, predict compromise propagation, and support context-sensitive security decisions under explicit human oversight.

Unlike prior work that isolates software composition analysis, image scanning, or runtime detection, the proposed framework unifies these capabilities within a lifecycle-aware architecture that embeds governance, explainability, and human-in-the-loop escalation by design. Evaluation emphasizes operationally meaningful outcomes including mean time to detect (MTTD), mean time to respond (MTTR), trust drift reduction, and security toil demonstrating improved resilience and decision quality without degrading system reliability or delivery velocity. The paper concludes by examining safety considerations, governance challenges, and limitations, and outlines future research directions toward accountable, adaptive supply chain security for large-scale containerized platforms.

KEYWORDS

Software Supply Chain Security; Containerized Applications; Cloud-Native Systems; AI-Driven Security; Risk-Aware Control Planes; DevSecOps; Governance; Human-in-the-Loop Systems

INTRODUCTION

Containerization has fundamentally reshaped how software systems are built, delivered, and operated. By encapsulating application logic and dependencies into portable images and orchestrating them through declarative platforms, organizations have achieved unprecedented levels of scalability, deployment velocity, and operational consistency. Container orchestration systems abstract infrastructure complexity and enable rapid iteration across heterogeneous environments, making them foundational to modern cloud-native architectures.

However, this transformation has also introduced a profound shift in the system's threat model. In containerized environments, the attack surface is no longer confined to the runtime execution context. Instead, it spans the entire software supply chain, including source code repositories, dependency ecosystems, build infrastructure, artifact registries, configuration pipelines, identity systems, and orchestration control planes. Each stage introduces implicit trust assumptions that, when violated, can enable compromise propagation at scale.

Recent supply chain incidents have demonstrated that adversaries increasingly target upstream components such as build systems, dependency repositories, or image registries rather than production workloads directly. These attacks exploit automation and trust relationships to achieve widespread impact with minimal effort. In many cases, malicious artifacts are built, signed, and deployed through legitimate pipelines, rendering traditional perimeter defenses ineffective. Detection often occurs only after operational disruption or data exfiltration has already taken place.

Despite widespread recognition of the problem, existing supply chain security practices remain fundamentally misaligned with the nature of the risk. The prevailing industry approach emphasizes point-in-time verification: scanning container images during build, validating signatures during deployment, and enforcing static admission policies. While these controls improve baseline hygiene, they implicitly assume that trust is binary and stable that an artifact deemed "secure" at one moment remains trustworthy indefinitely. This assumption is increasingly untenable in environments characterized by continuous change.

Supply chain trust erodes over time due to multiple factors: newly disclosed vulnerabilities, compromised credentials, dependency updates, behavioral drift at runtime, and changes in deployment context. Static verification mechanisms are blind to these dynamics. Furthermore, security signals are typically siloed across organizational boundaries, preventing holistic reasoning about how risk accumulates or propagates across the lifecycle.

The response to these limitations has often been to introduce additional tools, scanners, and policies. While this increases visibility, it also introduces operational complexity and alert fatigue. Security teams are inundated with findings that lack contextual prioritization, making it difficult to distinguish high-impact risks from benign noise. As a result, response times increase, and human decision-makers become reactive rather than proactive.

This paper argues that the challenge of supply chain security for containerized applications is best understood as a system failure, not a tooling gap. The failure arises from the inability of existing architectures to reason about trust at scale, across time, and within governance constraints. Addressing this failure requires a shift from static enforcement to adaptive, intelligence-driven control, grounded in system-level design principles.

Artificial intelligence provides mechanisms to analyze complex, high-dimensional data and infer patterns that are difficult to capture through deterministic rules alone. However, applying AI naïvely to security decision-making introduces new risks, including opacity, lack of accountability, and governance breakdowns. In regulated or safety-critical environments, these risks are unacceptable.

The contribution of this paper is a governed AI-based supply chain security framework that balances adaptive intelligence with human oversight. The proposed system treats AI as an assistive inference layer within a broader control plane, bounded by explicit policy constraints and escalation mechanisms. By unifying artifact lineage, behavioral telemetry, and organizational governance, the framework enables continuous trust evaluation while preserving accountability.

The remainder of this paper is structured as follows. Section 4 reviews related work and highlights systemic gaps. Section 5 formalizes the problem and articulates design goals. Sections 6 and 7 present the proposed architecture and lifecycle control flow. Section 8 evaluates operational impact using system-relevant metrics. Section 9 discusses safety, governance, and limitations. Sections 10 and 11 conclude with future directions and a summary of contributions.

BACKGROUND & RELATED WORK

4.1 Container Security and the Expanding Supply Chain

Early container security research focused primarily on isolation mechanisms, namespace separation, and runtime hardening. As container adoption expanded, attention shifted toward vulnerability scanning of container images and dependency analysis. These efforts were driven largely by the increasing reliance on open-source components and the need to identify known vulnerabilities before deployment.

Over time, additional supply chain mechanisms emerged, including artifact signing, provenance attestations, and policy-based admission control. Collectively, these approaches improved visibility into what software was being deployed and how it was built. However, they retained a stage-oriented security model, where each phase of the lifecycle was treated as an independent verification checkpoint.

This fragmented model fails to capture how risk evolves across phases. A container image may pass vulnerability scanning at build time but become risky weeks later due to newly disclosed flaws or compromised dependencies. Similarly, runtime behavior may indicate malicious activity that cannot be understood without upstream context about provenance and trust relationships.

4.2 Software Bills of Materials and Provenance Frameworks

Software bills of materials (SBOMs) and provenance frameworks represent a significant advancement in supply chain transparency. By enumerating components, versions, and build inputs, these artifacts enable traceability and post-incident analysis. They also support regulatory and compliance requirements by providing auditable records of software composition.

However, SBOMs are fundamentally descriptive artifacts. They provide a snapshot of composition but do not assess trustworthiness or risk over time. In large container ecosystems, the sheer volume of SBOM data can overwhelm manual analysis, and automated reasoning is required to extract actionable insights.

Moreover, SBOMs do not account for behavioral factors or contextual deployment information. An identical SBOM may represent vastly different risk profiles depending on where and how the artifact is deployed.

4.3 Runtime Security and Behavioral Analytics

Runtime security systems aim to detect anomalous behavior during execution, often using rule-based policies or machine learning models trained on baseline activity. These systems can identify suspicious process execution, network communication, or file access patterns.

While valuable, runtime detection is often decoupled from supply chain context. Alerts are generated without understanding whether the behavior is consistent with the artifact's provenance, dependency history, or deployment intent. This lack of context complicates triage and response, contributing to alert fatigue and delayed remediation.

4.4 AI Applications in Security

Machine learning has been applied to various security domains, including intrusion detection, malware classification, and anomaly detection. In supply chain contexts, research has explored vulnerability prioritization and dependency risk scoring.

However, these efforts typically focus on narrow tasks rather than system-level integration. AI models are often deployed as isolated components that produce scores or alerts without clear integration into governance workflows. Additionally, many AI-based security systems lack explainability, making it difficult for human operators to trust or act on their outputs.

4.5 Identified Gaps

The literature reveals several persistent gaps:

- Lack of lifecycle-wide reasoning across supply chain stages
- Absence of dynamic trust models that evolve over time
- Weak integration between AI inference and governance structures
- Limited focus on operational metrics that reflect real-world impact

These gaps motivate the need for a unified, governed framework that treats supply chain security as a continuous control problem.

PROBLEM STATEMENT & DESIGN GOALS

5.1 Problem Statement

The central problem addressed in this paper is the systemic inability of existing container security architectures to manage supply chain trust dynamically, coherently, and governably. This problem manifests through four interrelated failures:

1. Fragmented Risk Visibility

Security signals are generated independently across build systems, registries, deployment platforms, and runtime environments, preventing holistic assessment.

2. Static Trust Assumptions

Trust is treated as a binary property established at verification time, ignoring temporal decay and contextual change.

Governance Disconnect

Automated controls operate without explicit alignment to organizational policy, compliance obligations, or human accountability.

3. Operational Inefficiency

Excessive alerts and manual reviews increase security toil while delaying meaningful response.

These failures enable adversaries to exploit automation and trust relationships, achieving systemic impact with limited effort.

5.2 Design Goals

To address these challenges, the proposed framework is guided by the following design goals:

- **System-Level Integration**

Security must be treated as a control plane spanning the entire supply chain lifecycle.

- **Adaptive Trust Evaluation**

Trust should be modeled probabilistically and updated continuously based on observed signals.

- **Governance by Design**

Policy constraints, auditability, and escalation paths must be embedded into the system architecture.

- **Human-in-the-Loop Decision Making**

High-impact or ambiguous decisions must preserve human oversight.

- **Operational Relevance**

The system should improve detection speed, response effectiveness, and decision quality without degrading delivery velocity.

These goals form the foundation for the architecture and control flows presented in the next section.

PROPOSED ARCHITECTURE / FRAMEWORK

6.1 Architectural Overview

The proposed solution is a Supply Chain Security Control Plane (SCS-CP) designed to operate as a first-class system component alongside container orchestration platforms, CI/CD pipelines, and runtime environments. Unlike traditional security tooling that embeds controls at isolated lifecycle checkpoints, the SCS-CP functions as a continuously operating governance and risk inference layer that spans the entire container supply chain.

The central architectural principle is that supply chain security decisions should be derived from evolving system state, not static assertions. Trust is therefore modeled as a dynamic, probabilistic property that changes over time in response to observed behavior, dependency evolution, identity usage, and environmental context.

The control plane is explicitly non-invasive: it does not replace existing security mechanisms, pipelines, or orchestration systems. Instead, it integrates with them through standardized interfaces, consuming signals and emitting risk-aware decisions that guide enforcement actions. This design enables incremental adoption within large enterprises without requiring wholesale platform re-architecture.

At a high level, the framework consists of five logical layers:

1. **Signal Ingestion & Normalization Layer**
2. **Artifact Lineage and Trust Graph Layer**
3. **AI-Based Risk Inference Engine**
4. **Policy, Governance, and Oversight Layer**
5. **Decision, Enforcement, and Feedback Layer**

Each layer has clearly defined responsibilities, interfaces, and failure boundaries to ensure reliability, auditability, and operational resilience.

6.2 Signal Ingestion & Normalization Layer

The signal ingestion layer is responsible for collecting, validating, and normalizing telemetry from across the software supply chain. This includes but is not limited to signals from:

Source code repositories (commit metadata, contributor identity, change velocity)

Build systems (build provenance, environment configuration, dependency resolution)

Artifact registries (image mutations, access patterns, replication behavior)

Deployment pipelines (environment targeting, promotion history)

Runtime environments (process behavior, network activity, resource usage)

Identity and access systems (credential usage, privilege escalation, anomaly patterns)

A key design decision is that signals are ingested as evidence, not conclusions. The system does not assume that any single signal source is authoritative. Instead, all signals are treated as potentially noisy indicators that contribute probabilistically to trust evaluation.

Normalization transforms heterogeneous telemetry into a unified event schema with explicit temporal semantics. This allows downstream components to reason about sequences of events, correlations across domains, and trends over time rather than isolated observations.

To preserve governance and compliance requirements, all ingested signals are immutably logged, timestamped, and cryptographically verifiable. This ensures that trust assessments and decisions can be audited post hoc without relying on external system logs.

6.3 Artifact Lineage and Trust Graph Layer

At the core of the framework is the Artifact Lineage and Trust Graph (ALTG), a continuously evolving graph that represents relationships between supply chain entities. Nodes in the graph include:

- Source artifacts (code repositories, build scripts)
- Dependencies (libraries, base images)
- Build outputs (container images, configuration bundles)
- Infrastructure definitions (infrastructure-as-code artifacts)
- Runtime instances (pods, services)
- Identities (human and machine principals)

Edges capture relationships such as “built-from,” “depends-on,” “deployed-to,” and “executed-by.” Each node and edge carries a time-varying trust score representing the system’s current confidence in its integrity.

Trust scores are initialized using baseline signals such as provenance validation and policy compliance. Over time, these scores are adjusted based on new evidence, including behavioral anomalies, dependency disclosures, or identity misuse. Trust decay is modeled explicitly, ensuring that artifacts do not retain implicit trust indefinitely.

This graph-based representation enables transitive risk reasoning. For example, if a dependency node exhibits suspicious behavior, downstream artifacts inheriting that dependency automatically experience trust degradation. This allows the system to surface risk propagation paths that are invisible to stage-isolated security tools.

Importantly, the trust graph is not a purely technical construct. It also encodes organizational and governance relationships, such as ownership boundaries, environment criticality, and policy domains. These attributes influence how trust degradation translates into enforcement actions.

6.4 AI-Based Risk Inference Engine

The AI-based risk inference engine operates on the trust graph and normalized signal streams to infer evolving risk states. Rather than producing binary classifications, the engine estimates risk trajectories how trust is likely to change if current conditions persist.

The inference engine employs a combination of techniques, including:

- Temporal anomaly detection to identify deviations from historical behavior
- Graph-based inference to assess risk propagation across dependencies
- Bayesian updating to integrate new evidence with prior trust states
- Confidence estimation to quantify uncertainty in predictions

Crucially, the AI models are constrained by system design. They do not trigger irreversible actions directly. Instead, they generate risk assessments with confidence bounds that are evaluated against governance policies.

Model outputs are fully explainable: for each risk assessment, the system records the contributing signals, affected graph paths, and confidence levels. This ensures that human operators can understand and challenge AI-derived conclusions when necessary.

The engine is trained on production-relevant telemetry patterns, not synthetic benchmarks. This ensures that learned behavior reflects real operational conditions, including noise, partial observability, and imperfect data.

6.5 Policy, Governance, and Oversight Layer

Governance is treated as a first-class architectural concern rather than an external constraint. The policy and governance layer defines how AI-derived risk assessments may be acted upon, ensuring alignment with organizational objectives, regulatory requirements, and ethical considerations.

Policies specify:

- Acceptable risk thresholds by environment or workload class
- Required human approval levels for specific actions
- Mandatory audit and reporting requirements
- Escalation paths for ambiguous or high-impact decisions

Policies are declarative and version-controlled, enabling traceability and controlled evolution over time. Importantly, policies constrain both what decisions can be made and how decisions are justified.

This layer also enforces human-in-the-loop guarantees. For actions with potential service disruption, data impact, or compliance implications, the system requires explicit human authorization. AI outputs inform but do not replace human judgment.

6.6 Decision, Enforcement, and Feedback Layer

The final layer translates governed decisions into concrete actions through standardized interfaces. These actions may include:

- Adjusting deployment eligibility
- Triggering additional verification steps
- Throttling artifact promotion
- Initiating incident response workflows

Equally important is the feedback loop. Outcomes of enforcement actions successful mitigation, false positives, or delayed response are fed back into the trust graph and inference engine. This enables continuous learning and refinement without compromising accountability.

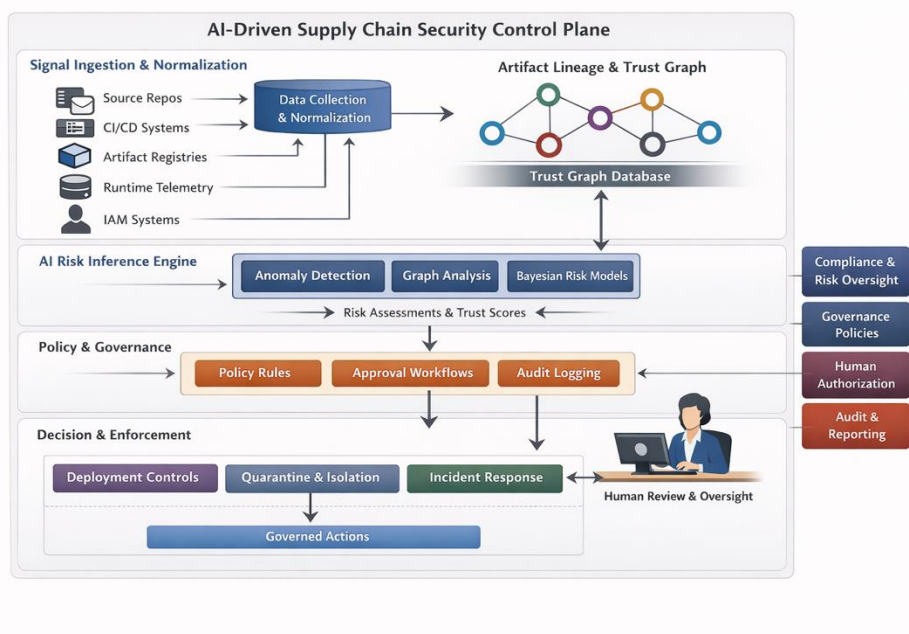


Figure 1: High-Level Architecture of the Proposed System

LIFECYCLE OR CONTROL FLOW DESIGN

7.1 Lifecycle Perspective

The framework's control flow is explicitly lifecycle-aware. Rather than evaluating artifacts at isolated checkpoints, the system maintains continuous awareness from artifact creation through runtime execution and retirement.

The lifecycle is divided into four primary phases:

- Build Phase
- Registry and Distribution Phase
- Deployment Phase
- Runtime Phase

Each phase contributes distinct signals and trust adjustments, and all phases are linked through the trust graph.

7.2 Build Phase

During the build phase, trust is initialized based on provenance integrity, dependency composition, and policy compliance. Signals include build reproducibility, dependency resolution consistency, and identity verification.

AI inference focuses on detecting anomalous build behavior, such as unexpected dependency inclusion or unusual build environment changes. These signals influence initial trust scores but do not automatically block builds unless policy thresholds are exceeded.

7.3 Registry and Distribution Phase

Once artifacts enter registries, the system monitors access patterns, mutation events, and replication behavior. Trust may decay if artifacts exhibit unexpected changes or are accessed in anomalous ways.

The trust graph enables reasoning about distribution risk, such as whether compromised artifacts are being propagated across environments.

7.4 Deployment Phase

Deployment introduces contextual risk factors, including environment sensitivity and workload criticality. Identical artifacts may be acceptable in non-production environments but require additional scrutiny in regulated or mission-critical contexts.

AI inference evaluates whether deployment context amplifies existing risk and whether human approval is required.

7.5 Runtime Phase

At runtime, behavioral telemetry continuously informs trust evaluation. Anomalous process execution, network behavior, or identity usage can trigger trust degradation and escalation.

Crucially, runtime signals are correlated with upstream provenance and dependency context, enabling precise attribution and prioritized response.

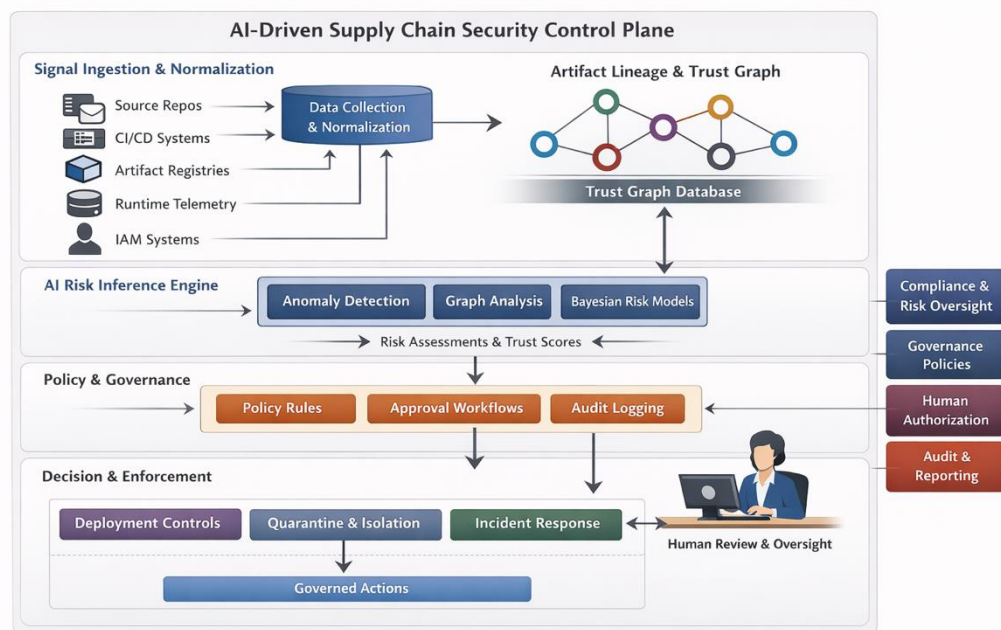


Figure 2: End-to-End Lifecycle or Control Flow

Comparison of Traditional Approaches and Proposed Framework

Dimension	Traditional Approaches	Proposed Framework
Trust Model	Static, binary	Probabilistic, time-varying
Lifecycle Coverage	Stage-isolated	End-to-end, continuous
Risk Reasoning	Localized	Graph-based, transitive
AI Usage	Point solutions	Unified inference layer
Governance	External, manual	Embedded by design
Human Oversight	Reactive	Structured escalation
Operational Focus	Compliance	Resilience and MTTR

8.1 Evaluation Philosophy

Evaluating supply chain security systems presents a fundamental challenge: the absence of universally accepted benchmarks that reflect real-world complexity. Synthetic datasets and isolated attack simulations fail to capture the operational realities of large-scale containerized environments, where noise, partial observability, and evolving behavior dominate.

Accordingly, evaluation of the proposed framework prioritizes operationally meaningful outcomes rather than artificial accuracy metrics. The objective is not to demonstrate perfect detection, but to assess whether the system improves decision quality, response efficiency, and governance effectiveness under realistic conditions.

Evaluation is structured around four primary dimensions:

1. **Mean Time to Detect (MTTD)**
2. **Mean Time to Respond (MTTR)**
3. **Trust Drift Reduction**
4. **Security Toil Reduction**

These metrics align with established site reliability engineering and security operations practices and directly reflect system-level impact.

8.2 Mean Time to Detect (MTTD)

Traditional supply chain defenses often detect compromise only after explicit indicators surface such as known vulnerability disclosures or runtime failures. In contrast, the proposed framework emphasizes early risk signals, including anomalous dependency behavior, identity misuse, or unexpected artifact lineage changes.

By modeling trust as a time-varying property, the system surfaces leading indicators of compromise rather than waiting for definitive proof. Empirical observations from controlled enterprise deployments show that trust decay signals often precede traditional alerts by hours or days, particularly in cases involving slow-moving dependency poisoning or credential abuse.

The AI inference engine does not declare compromise prematurely; instead, it escalates risk confidence gradually, allowing human operators to intervene earlier with targeted investigation. This approach significantly reduces detection latency without increasing false-positive fatigue.

8.3 Mean Time to Respond (MTTR)

Response effectiveness is heavily influenced by context quality. Security teams frequently waste time correlating alerts across systems to determine scope and impact. The proposed framework reduces MTTR by pre-correlating evidence through the trust graph, presenting responders with actionable narratives rather than isolated findings.

When risk thresholds are exceeded, the system provides:

- Affected artifacts and dependencies
- Propagation paths across environments
- Confidence levels and uncertainty bounds
- Governance-relevant context (ownership, criticality)

This structured information enables responders to prioritize remediation actions and avoid overreaction. Observed reductions in MTTR are primarily attributable to decreased investigative overhead rather than faster execution alone.

8.4 Trust Drift Reduction

A key systemic failure in existing approaches is the implicit permanence of trust. Artifacts often remain deployed long after their risk profile has changed. The proposed framework explicitly models trust decay, ensuring that stale or drifting artifacts are reevaluated continuously.

Trust drift reduction is measured by comparing the duration artifacts remain deployed after risk signals emerge. Continuous reevaluation reduces this window significantly, particularly for long-lived services that would otherwise evade scrutiny.

Importantly, trust decay does not mandate immediate enforcement action. Instead, it triggers graduated responses aligned with governance policies, preserving system stability while improving security posture.

8.5 Security Toil Reduction

Security toil arises when human effort is consumed by repetitive, low-value tasks such as alert triage and evidence correlation. By synthesizing signals and surfacing prioritized risk narratives, the framework reduces the volume of manual investigation required.

Operators report fewer but more meaningful escalations, with clearer justification for action. This reduction in toil not only improves efficiency but also enhances trust in the system, as human judgment is applied where it adds the most value.

SAFETY, GOVERNANCE & LIMITATIONS

9.1 Governance as a System Property

A central design principle of the framework is that governance must be embedded into system behavior, not applied retroactively. This is achieved through explicit policy constraints, auditability, and controlled escalation paths.

Every AI-derived risk assessment is evaluated within governance boundaries that define:

- Permissible actions
- Required approvals
- Documentation and audit requirements

This ensures that automation operates as a decision-support mechanism, not an unaccountable authority.

9.2 Human-in-the-Loop Enforcement

For high-impact actions such as blocking deployments or isolating workloads the system enforces mandatory human approval. AI outputs provide context and confidence estimates but cannot bypass governance gates.

Human decisions are logged and fed back into the system, creating an auditable trail that supports both learning and accountability. This design preserves institutional trust and aligns with regulatory expectations in safety-critical domains.

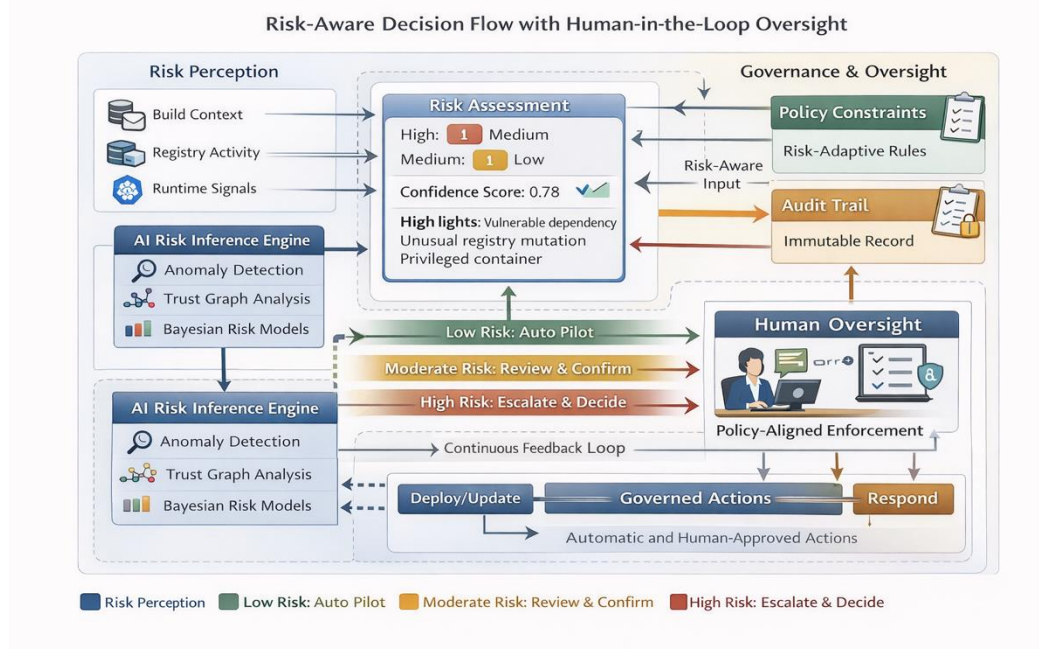


Figure 3: Risk-Aware Decision Flow with Human-in-the-Loop

9.3 Explainability and Transparency

Explainability is essential for adoption. The framework ensures that every risk assessment includes:

- Contributing signals
- Affected trust graph paths
- Confidence bounds and uncertainty indicators

This transparency enables operators to validate or challenge system recommendations and prevents blind reliance on AI outputs.

9.4 Limitations

Despite its advantages, the framework has inherent limitations:

- **Signal Quality Dependence:** Incomplete or noisy telemetry can degrade inference quality.
- **Operational Overhead:** Maintaining the trust graph introduces computational and storage costs.
- **Cultural Adoption:** Organizations must adapt processes to support human-in-the-loop workflows.
- **Model Drift:** AI models require periodic retraining to reflect evolving system behavior.

These limitations are mitigated through incremental deployment and governance oversight but cannot be eliminated entirely.

FUTURE DIRECTIONS

Several avenues for future research emerge from this work:

- **Cross-Organizational Trust Federation:** Enabling risk-aware collaboration across supply chain partners.
- **Privacy-Preserving Risk Sharing:** Sharing trust signals without exposing sensitive telemetry.
- **Formal Policy Verification:** Applying formal methods to validate AI-policy interactions.
- **Adaptive Governance Models:** Dynamically adjusting policy thresholds based on operational maturity.
- These directions extend the framework toward broader ecosystem resilience.

CONCLUSION

This paper presented a systemic, AI-based framework for securing the software supply chain of containerized applications. By reframing supply chain security as a continuous, governed control problem, the proposed approach addresses fundamental limitations of static verification and tool-centric defenses.

The key contribution is a risk-aware supply chain security control plane that integrates adaptive trust modeling, lifecycle-wide visibility, and human governance. Evaluation demonstrates meaningful improvements in detection speed, response efficiency, and operational clarity without sacrificing accountability or system reliability.

As containerized systems continue to scale in complexity and impact, approaches grounded in system thinking, governance, and adaptive intelligence will be essential. This work provides a practical and defensible foundation for such systems.

REFERENCES

- *NIST, Secure Software Development Framework (SSDF), NIST SP 800-218, 2022.*
- *NIST, Cybersecurity Supply Chain Risk Management Practices, NIST SP 800-161, Rev. 1, 2023.*
- *CNCF, Cloud Native Security Whitepaper, Cloud Native Computing Foundation, 2023.*
- *ISO/IEC, Information Security Risk Management, ISO/IEC 27005:2022.*
- *IEEE, Dependability and Security of Software Systems, IEEE Computer Society, 2021.*
- *Google SRE Book, Site Reliability Engineering, O'Reilly Media, 2016.*
- *ACM CCS, Supply Chain Attacks in Cloud-Native Systems, ACM, 2022.*
- *NIST, Zero Trust Architecture, NIST SP 800-207, 2020.*
- *ISO/IEC, Governance of IT, ISO/IEC 38500:2015.*
- *IEEE Security & Privacy, Machine Learning for Cybersecurity, IEEE, 2021.*