



Tx Cloud and AI Security Practice

WWW.TESTINGXPERTS.COM





Cloud and AI Security



Cloud Security Configuration Audit

- Follows CIS Benchmarks and Organizations Security Baselines.
- Audit cloud security configuration against baseline or standards to identify security weaknesses or loopholes.
- Automated and Manual Processes to ensure complete coverage and efficient testing.
- Coverage across Amazon Web Services, Google Cloud Platform, Azure, M365, Oracle Cloud, Google Workspace etc.



Managed Cloud Security

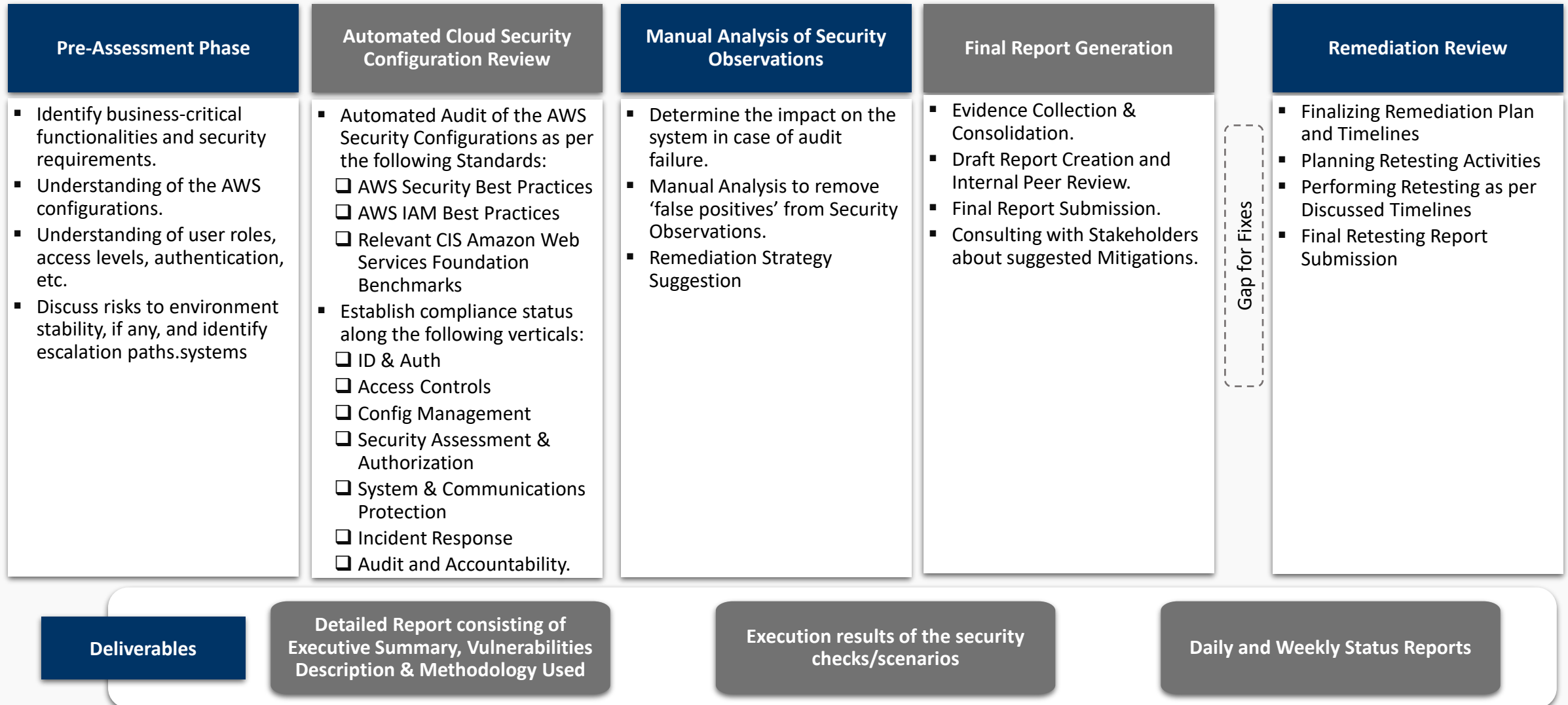
- Coverage for AWS, Google GCP, Azure, M365 and Google Workspace.
- Partner with clients to take ownership of managing security configuration across their cloud platform.
- Shared Responsibility Model – We manage security, Clients manage functionality, Cloud Provider manages availability.



Cloud Governance, Risk & Compliance

- Coverage for AWS, Google GCP, Azure, M365 and Google Workspace.
- Provide advisory and consultation to clients regarding formulation and execution of effective and compliant Cloud Security policies and procedures.
- Enabling Cloud Security Monitoring and Audit Logging to assist in effective Governance and Risk Management.

Comprehensive Cloud Security Configuration Audit Approach





AI / LLM Application Security Testing

- Follows latest industry standards in AI / LLM Security Testing.
- Identify Security Risks with an 'Outside-In' and 'Inside-Out' testing approach.
- Assures Robust Implementation of LLM's within applications.
- Works for offline and online LLM models.
- Coverage for OWASP Top 10 for LLM Applications 2024 & OWASP Top 10 for LLM Application 2025.
- Comprehensive checklists of AI/LLM Specific checks custom designed for each specific engagement needs.



AI / LLM Model Security Testing Coverage



- **LLM Prompt Injection:** Tests identify prompt manipulation that subverts model intent and control boundaries.



- **Sensitive Information Disclosure:** Ensures outputs don't leak confidential, personal, or regulated data.



- **Supply Chain Vulnerabilities:** Assesses risks from unverified models, datasets, and upstream dependencies.



- **Data and Model Poisoning:** Detects tampered training inputs that degrade accuracy or inject bias.



- **Improper Output Handling:** Validates that outputs are safely interpreted, filtered, and constrained.



- **Excessive Agency:** Flags over-permissioned models capable of unsafe or autonomous actions.



- **System Prompt Leakage:** Prevents exposure of hidden instructions that could be reverse-engineered.



- **Vector and Embedding Weakness:** Probes embedding logic for tampering, leakage, or semantic collisions.



- **Misinformation:** Evaluates risks of confidently generated but factually incorrect content.



- **Unbounded Consumption:** Tests for limits on input volume and frequency to prevent resource exhaustion.

A background image showing a close-up of two hands shaking in a firm grip, symbolizing a business deal or agreement. The image is dark and moody, with a large, bright red diagonal shape overlaying the right side. The text 'Thank you' is centered in a white box on the left side.

Thank you