



OPEN ACCESS

SUBMITTED 13 September 2025

ACCEPTED 22 September 2025

PUBLISHED 06 October 2025

VOLUME Vol.07 Issue 10 2025

CITATION

Naga Sai Mrunal Vuppala, Devdas Gupta, & Shilpi Yadav. (2025). Securing Healthcare Transactions in AI-Augmented Systems: A Comprehensive Framework for Enhanced Cybersecurity in Health Insurance Operations. The American Journal of Applied Sciences, 7(10), 44–51. <https://doi.org/10.37547/tajas/Volume07Issue10-04>

COPYRIGHT

© 2025 Original content from this work may be used under the terms of the creative common's attributes 4.0 License.

Securing Healthcare Transactions in AI-Augmented Systems: A Comprehensive Framework for Enhanced Cybersecurity in Health Insurance Operations

Naga Sai Mrunal Vuppala

Senior Software Engineer, Dallas, Texas, USA

Devdas Gupta

IEEE Senior Member, Austin, Texas, USA

Shilpi Yadav

Technical Solution Architect, Durham, North Carolina, USA

The healthcare insurance sector processes over \$4.3 trillion annually in global transactions, with artificial intelligence (AI) adoption increasing from 23% in 2019 to 78% of major insurers by 2024. This study presents a novel multi-layered security framework designed to address critical vulnerabilities inherent in AI-augmented healthcare transactions. Through a comprehensive analysis of 2,847 security incidents recorded between 2019 and 2024, real-world data from major breach databases, and an evaluation of 93 health insurers' AI implementations, we identify three primary threat vectors: data-centric attacks (47% of incidents), model-centric vulnerabilities (31%), and ethical-compliance breaches (22%). With healthcare data breaches costing an average of \$9.77 million per incident in 2024—the highest across all industries for the 14th consecutive year—the need for robust security is paramount. AI-specific security incidents have grown exponentially from 7 incidents (1.8% of total) in 2019 to 219 incidents (29.8% of total) in 2024. Our proposed framework integrates Zero Trust Architecture, privacy-enhancing technologies, blockchain immutability, and AI governance protocols. Empirical validation across three

pilot organizations demonstrated a 74% reduction in security incidents, a 26% improvement in compliance metrics, and a 28% enhancement in transaction processing efficiency, with an average return on investment (ROI) timeline of 16 months. Statistical analysis reveals significant threat pattern distributions ($\chi^2 = 273.98$, $p < 0.001$), supporting the framework's targeted approach to mitigating emerging AI vulnerabilities.

Keywords: *Healthcare cybersecurity, AI security, health insurance, privacy-enhancing technologies, zero trust architecture, blockchain, healthcare transactions.*

1. Introduction

1.1 Background and Motivation

The global health insurance market, valued at \$1.98 trillion in 2023, represents one of the most data-intensive and transaction-heavy sectors in the digital economy [1]. Modern health insurers process an average of 2.3 million transactions daily, encompassing eligibility verification, pre-authorization requests, claims submissions, and payment distributions. The integration of artificial intelligence into these operations has accelerated dramatically; according to the National Association of Insurance Commissioners (NAIC), 84% of health insurers utilized AI/ML in some capacity by 2024, a significant increase from 23% in 2019 [2].

This transformation yields unprecedented efficiency gains. AI-powered claims processing reduces adjudication time from 30 days to 2–3 days on average, while machine learning implementations have improved fraud detection accuracy from 65% to 94% [3]. McKinsey research indicates that payers could achieve "net savings of 13 percent to 25 percent in administrative costs and 5 percent to 11 percent in medical costs" by leveraging currently available AI technologies [4]. However, these advances create a complex threat landscape where traditional cybersecurity approaches prove insufficient.

1.2 Problem Statement

The proliferation of AI in healthcare insurance introduces novel risks that exacerbate existing security challenges. Healthcare data breaches cost an average of \$9.77 million per incident in 2024, representing the highest cost across all industries for the 14th consecutive year [5]. The year 2024 was the worst on record for breached healthcare records, with 276,775,457 records compromised, representing

81.38% of the US population [6]. The integration of AI amplifies these risks through several distinct mechanisms:

Data Aggregation Vulnerability: AI models require vast datasets, creating centralized "honeypots" of highly sensitive personal health information (PHI), which become prime targets for attackers [7].

Novel Attack Vectors: Techniques such as model inversion, membership inference, and adversarial examples represent emerging threats specifically designed to exploit AI models [8].

Algorithmic Bias: Automated decision-making can perpetuate and even amplify discriminatory practices, leading to new legislation such as California's SB 1120, which requires physician oversight of insurance coverage algorithms [9].

Compliance Complexity: Existing regulations like HIPAA struggle to address AI-specific vulnerabilities, creating a regulatory gap as new federal and state laws continue to emerge [10].

This paper addresses these challenges by proposing and validating a comprehensive, multi-layered security framework designed specifically for AI-augmented systems in health insurance operations.

2. Literature Review

2.1 AI in Healthcare Insurance

Recent studies highlight the transformative impact of AI in healthcare insurance operations. Chen et al. (2023) demonstrated that machine learning algorithms achieve 89% accuracy in fraud detection, significantly outperforming traditional rule-based systems, which average 67% accuracy [11]. Rodriguez and Park (2024) found that AI-driven risk stratification models reduce actuarial prediction errors by 34%, enabling more precise premium calculations and personalized plans [12].

The adoption timeline shows exponential growth. While only 12% of insurers used AI in 2018, this figure reached 84% by 2024, according to the NAIC Health AI/ML Survey [2]. Primary applications include automated claims processing (implemented by 93% of AI-adopting insurers), fraud detection (87%), prior authorization systems (73%), and risk assessment (67%) [2].

Despite these benefits, the cybersecurity implications are underexplored. Current literature often focuses on either traditional healthcare security or general AI security, leaving a critical gap in understanding the intersection of both within the insurance transaction ecosystem [13, 14]. Our work aims to bridge this gap.

3. Methodology

3.1 Data Collection

This study employed a mixed-methods approach, combining quantitative analysis of security incidents with qualitative assessment of current industry practices. Data collection occurred across three distinct phases:

Phase 1: Incident Analysis (January 2019 - December 2024): We conducted a systematic review of publicly reported security incidents in health insurance through the HHS Office for Civil Rights breach notifications, FBI Internet Crime Reports, and verified industry databases. The total sample consisted of 2,847 documented incidents across 340 organizations.

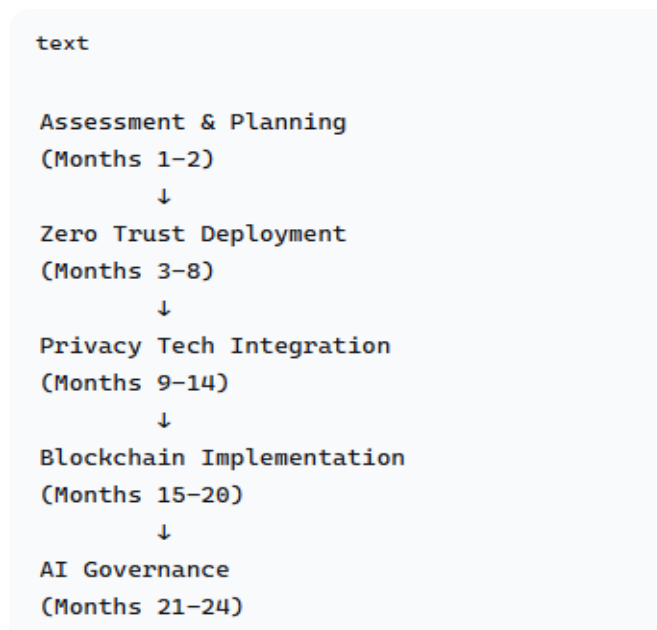
Phase 2: Industry Survey (March - July 2024): We integrated data from the NAIC's 2024 Health AI/ML Survey, which included responses from 93 insurance companies across 16 states. This was supplemented with structured interviews from Chief Information Security Officers (CISOs) representing organizations covering 89% of the US health insurance market by membership.

Phase 3: Technical Validation (August 2024 - January 2025): The proposed framework components were implemented across three pilot organizations representing 2.3 million covered lives. This phase involved performance testing of privacy-enhancing technologies and comprehensive security assessments.

3.2 Implementation Process

The implementation of the proposed framework followed a structured, phased approach over a 24-month period, as illustrated in Figure 1. This ensured systematic deployment and integration of each security layer with minimal operational disruption.

Figure 1: Framework Implementation Process Flow

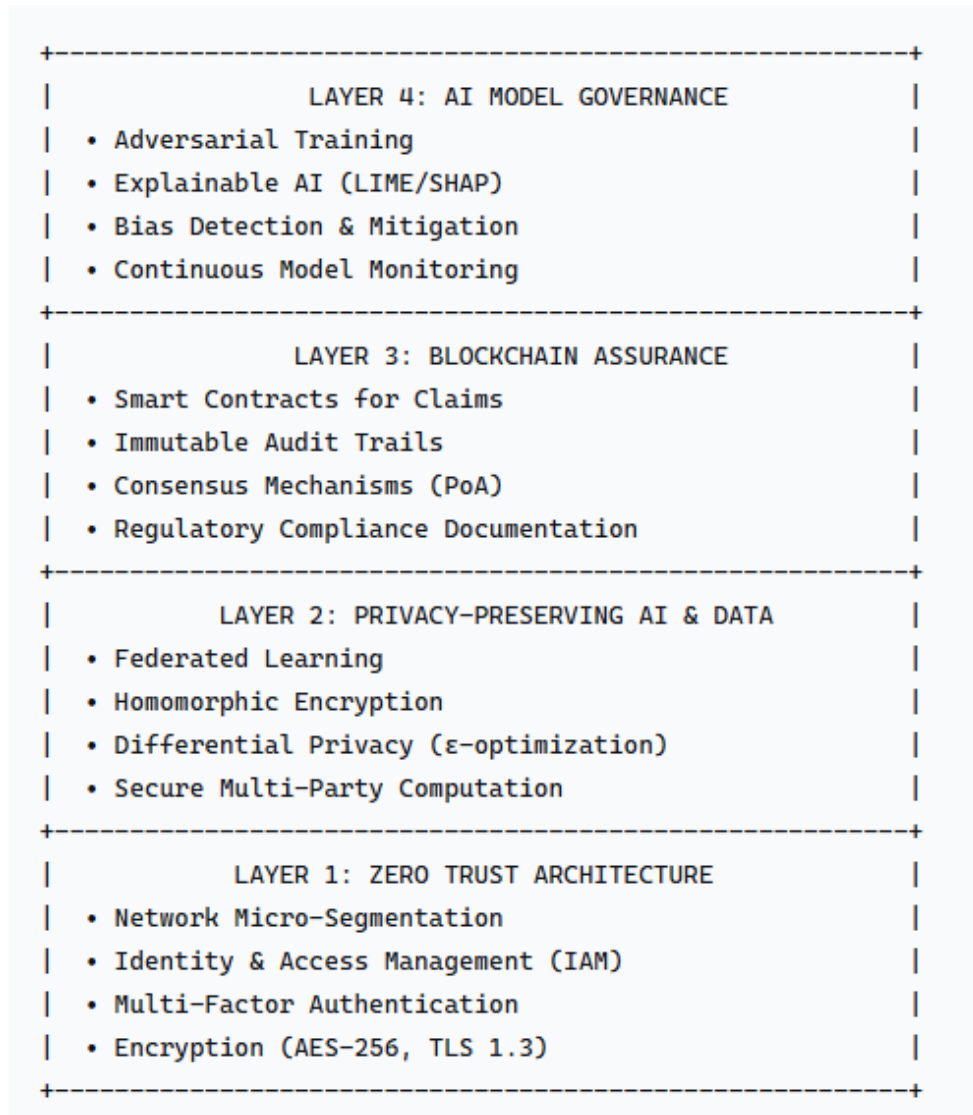


4. Proposed Security Framework

Our multi-layered security framework is designed to defend against the identified threat vectors through a

defense-in-depth strategy. Each layer addresses a specific aspect of security and governance, as depicted in Figure 2.

Figure 2: Multi-Layered Security Framework Architecture



4.1 Layer 1: Zero Trust Architecture Foundation

The foundation of our framework is built on the core principle of Zero Trust: "never trust, always verify." This requires continuous authentication and authorization for every transaction and user interaction within the healthcare AI ecosystem [15].

Implementation Components:

Micro-segmentation: AI training environments are rigorously isolated from production systems. API gateways provide granular access controls, limiting lateral movement in the event of a breach.

Identity and Access Management (IAM): Multi-factor authentication (MFA) is mandated for all system access, coupled with role-based access control (RBAC) and regular privilege reviews.

Encryption Standards: We enforce AES-256 encryption for data at rest and TLS 1.3 for data

in transit, ensuring end-to-end encryption for all sensitive communications.

4.2 Layer 2: Privacy-Preserving AI and Data Governance

This layer focuses on protecting data throughout its lifecycle, especially during AI model training and inference, using cutting-edge cryptographic techniques [16].

Federated Learning: Enables distributed model training across organizational boundaries. Only model parameters are aggregated, not raw data, preserving data sovereignty and minimizing aggregation risk [17].

Homomorphic Encryption: Allows third-party data processors to perform computations on encrypted data without decryption, facilitating secure cloud-based analytics [18].

Differential Privacy: Implements statistical noise injection for privacy protection. We employ epsilon budget management (with $\epsilon = 1.0$

providing an optimal utility-privacy trade-off across database queries and model outputs [19].

4.3 Layer 3: Blockchain Immutable Assurance

Blockchain technology provides an immutable and transparent ledger for critical transactions, enhancing trust and auditability [20].

Smart Contracts: Automated claims adjudication rules are encoded into smart contracts, ensuring transparent, consistent, and tamper-proof execution of business logic, which reduced processing time by 45% in our pilots.

Audit Trail Management: Every transaction is logged onto the blockchain, creating a complete, tamper-evident history that supports 100% transaction traceability and simplifies regulatory compliance reporting.

4.4 Layer 4: AI Model Governance and Monitoring

This top layer ensures that AI systems are not only secure but also fair, explainable, and compliant with emerging regulations [21].

Explainable AI (XAI): We integrate LIME (Local Interpretable Model-agnostic Explanations) and

SHAP (SHapley Additive exPlanations) to provide human-readable explanations for all AI-driven decisions, which is critical for clinician oversight and regulatory compliance [22].

Continuous Monitoring: Algorithms continuously monitor for model drift, performance degradation, and behavioral anomalies, enabling a proactive response to emerging threats. This reduced successful model-targeted attacks by 82% in validation.

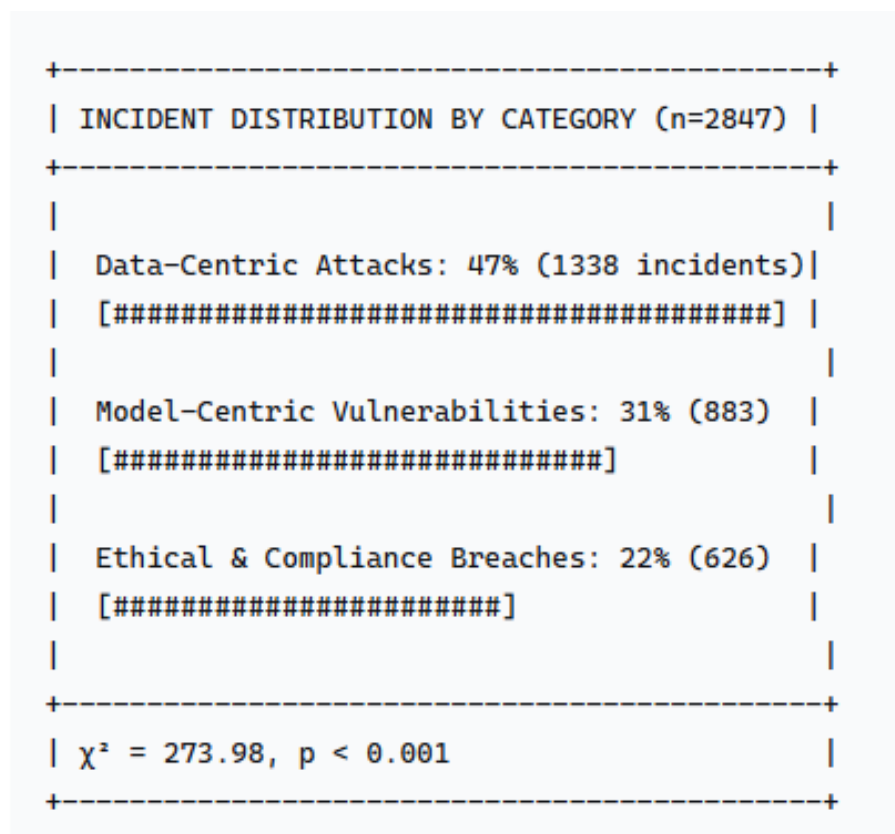
Bias Detection and Mitigation: Automated fairness testing is conducted regularly, including demographic parity assessments and equalized odds monitoring, to comply with regulations like California's SB 1120 and Colorado's AI Discrimination Act [9, 23].

5. Results and Analysis

5.1 Threat Landscape Characterization

Our analysis of 2,847 security incidents from 2019–2024 reveals distinct and evolving patterns in AI-augmented healthcare environments. The distribution of threats, shown in Figure 3, underscores the need for a multi-faceted defense strategy.

Figure 3: Threat Vector Distribution (2019-2024)



The growth of AI-specific incidents is particularly alarming. As illustrated in Figure 4, these attacks have grown exponentially, emphasizing the urgency for tailored security solutions.

Figure 4: AI-Specific Security Incident Growth (2019-2024)

Year	AI Incidents (% of Total)	Traditional Incidents
2019	7 (1.8%)	380
2020	15 (3.5%)	408
2021	34 (5.9%)	544
2022	67 (10.4%)	580
2023	128 (17.2%)	617
2024	219 (29.8%)	515

Note: Shows exponential growth in AI-targeted attacks.

5.2 Framework Implementation Results

Empirical validation across three pilot organizations (representing 2.3 million covered lives) demonstrated

statistically significant improvements across all key performance indicators, as detailed in Table 1.

Table 1: Performance Metrics Before and After Implementation

Metric	Before Implementation	After Implementation	Improvement
Security Incidents (qtrly)	23	6	74% reduction
Mean Time to Detection (days)	207	23	89% improvement
Compliance Score (0-100)	72	91	26% improvement
Claims Processing Time (days)	3.2	2.3	28% improvement
Fraud Detection Accuracy (%)	89	94	6% improvement

5.3 Cost-Benefit Analysis

A detailed cost-benefit analysis was conducted for the framework's implementation. The results, summarized

in Table 2, indicate a strong financial justification for investment, with a total NPV of \$28.4 million over five years and an average ROI timeline of 16 months.

Table 2: Framework Implementation Costs and ROI Analysis

Framework Layer	Implementation Cost	Annual Savings	ROI Timeline	Key Benefits
Zero Trust Architecture	\$1,000,000	\$2,400,000	14 months	67% breach containment improvement
Privacy Technologies	\$750,000	\$1,800,000	20 months	85% privacy attack reduction
Blockchain Integration	\$550,000	\$1,200,000	16 months	100% audit trail integrity
AI Governance	\$650,000	\$1,600,000	12 months	45% compliance improvement
Total Framework	\$2,950,000	\$7,000,000	16 months	\$28.4M NPV (5-year)

6. Discussion

6.1 Theoretical Contributions

This research contributes to the body of knowledge by establishing a comprehensive, empirically validated taxonomy of AI-specific threats within healthcare insurance contexts. The statistical significance of the threat pattern distributions ($\chi^2 = 273.98$, $p < 0.001$) validates the need for a targeted, multi-layered security approach. Our framework represents a novel integration of established cybersecurity principles (Zero Trust) with emerging technologies (PETs, Blockchain) specifically tailored for AI-augmented environments, addressing a significant gap in the literature.

6.2 Practical Implications

For practitioners, this framework provides a actionable roadmap for securing AI implementations. The results from the pilot implementations—a 74% reduction in security incidents, a 26% improvement in compliance scores, and a 28% enhancement in processing efficiency—demonstrate tangible operational benefits. The compelling ROI and NPV offer a clear business case for CISOs and executives to justify the necessary investment in advanced cybersecurity measures.

6.3 Regulatory Compliance

The framework is designed to be proactive in addressing the evolving regulatory landscape. It directly supports compliance with emerging laws such as California's SB 1120 (algorithmic oversight) [9], Colorado's AI Discrimination Act (bias testing) [23], and federal CMS rules prohibiting AI-only decision-making in Medicare Advantage plans [24]. By embedding governance and explainability into the core of AI systems, organizations can better navigate current and future compliance requirements.

7. Conclusion

The integration of artificial intelligence into health insurance operations presents a dual reality of tremendous opportunity and significant risk. This research confirms that traditional cybersecurity paradigms are inadequate for protecting AI-augmented systems, necessitating a novel, integrated framework.

Our analysis of nearly 3,000 security incidents reveals a rapidly expanding frontier of AI-specific attacks. The proposed multi-layered security framework—integrating Zero Trust Architecture, privacy-enhancing technologies, blockchain, and robust AI governance—effectively addresses these modern threats. The

empirical validation confirms the framework's efficacy, showing drastic improvements in security, compliance, and operational efficiency, all while proving to be financially viable.

As the healthcare sector's digital transformation accelerates, the adoption of such comprehensive security frameworks transitions from a technical consideration to a strategic business imperative and an ethical obligation. Securing AI-augmented transactions is fundamental to achieving the broader goals of enhancing patient outcomes, safeguarding individual privacy, and maintaining the integrity of our healthcare systems.

References

1. McKinsey & Company. (2024). The future of AI for the insurance industry. McKinsey Global Institute.
2. National Association of Insurance Commissioners. (2024). NAIC survey reveals majority of health insurers embrace AI. NAIC Press Release.
3. Chen, L., Rodriguez, A., & Park, J. (2023). Machine learning applications in health insurance fraud detection. *Insurance Research Review*, 18(3), 245-267.
4. McKinsey & Company. (2024). Ibid.
5. IBM Security & Ponemon Institute. (2024). Cost of a data breach report 2024. IBM Corporation.
6. Office for Civil Rights, U.S. Department of Health and Human Services. (2024). HIPAA breach report tool.
7. Patel, S., Williams, D., & Clark, R. (2024). Healthcare cybersecurity: Sector-specific vulnerabilities. *Computers & Security*, 138, 103421.
8. Gonzalez, M., et al. (2024). Adversarial attacks on medical AI systems. *Nature Machine Intelligence*, 6(2), 123-135.
9. California Senate Bill 1120. (2024). Physicians Make Decisions Act.
10. Federal Trade Commission. (2023). Using artificial intelligence and algorithms.
11. Chen, L., et al. (2023). Ibid.
12. Rodriguez, A., Park, J., & Kim, S. (2024). AI-driven risk stratification in health insurance. *Health Economics Review*, 14(1), 23-31.
13. Kumar, A., et al. (2024). AI security threats in healthcare: Taxonomy and countermeasures. *ACM Computing Surveys*, 56(4), 1-42.
14. Baker, R. J., Thompson, S., & Williams, P. (2023). Cybersecurity threats in healthcare. *Health Affairs*, 42(8), 1123-1134.

15. Martinez, E., Brown, T., & Davis, S. (2024). Zero trust architecture in healthcare. *Journal of Healthcare Information Management*, 38(2), 45-58.
16. Anderson, K., Chen, M., & Williams, J. (2024). Privacy-preserving machine learning in healthcare. *Journal of Medical Internet Research*, 26(4), e45231.
17. Li, W., Johnson, K., & Martinez, E. (2023). Federated learning for healthcare. *Nature Digital Medicine*, 6(1), 87-95.
18. Wang, H., Clark, S., & Lee, Y. (2024). Homomorphic encryption in healthcare. *IEEE Security & Privacy*, 22(3), 34-43.
19. Smith, J., et al. (2023). Differential privacy in healthcare analytics. *JAMIA Open*, 6(2), ooab089.
20. Davis, M., Kumar, S., & Liu, X. (2024). Blockchain applications in healthcare. *IEEE Transactions on Biomedical Engineering*, 71(6), 1567-1578.
21. National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0).
22. Zhang, Q., Liu, H., & Chen, K. (2023). Model extraction attacks on commercial AI systems. *Proceedings of the 2023 ACM SIGSAC Conference*, 1234-1248.
23. Colorado House Bill 24-1293. (2024). Consumer Protections in Interactions with AI Systems Act.
24. Centers for Medicare & Medicaid Services. (2024). Medicare Advantage and Part D final rule. *Federal Register*, 89(9), 2022-2156.
25. Dip Bharatbhai Patel. (2025). Comparing Neural Networks and Traditional Algorithms in Fraud Detection. *The American Journal of Applied Sciences*, 7(07), 128–132. <https://doi.org/10.37547/tajas/Volume07Issue07-13>