

Improvements of AI in the Analysis of Profitable Medical Insurance Strategie

Andreea-Miruna GHEORGHESCU

University of Economic Studies, Bucharest, Romania
gheorghescuandreea24@stud.ase.ro

Alexandru Gabriel ANDREI

University of Economic Studies, Bucharest, Romania
andreialexandru24@stud.ase.ro

Abstract:

In the complex healthcare system of the United States, adaptability through technological integration has become a primary driver of corporate longevity. This study explores how Artificial Intelligence (AI) mechanisms transform data into a strategic asset to secure a sustainable competitive advantage for medical insurance companies. The research evaluates the interplay between key demographic variables, including age, BMI, smoking status and insurance charges, in order to synthesize where AI can become a pivotal role in the company. Through both literature review and statistical analysis, it becomes obvious how the insurance field needs to react to the market quicker and better which entails the integration of algorithms. The results show that the transition from static data organization to dynamic, AI-enabled strategic optimization is the essential differentiator for companies seeking to lead the demographically complex U.S. health insurance market, since the improvements are consistent across the process of creating insurance plans for the consumers.

Keywords: AI, insurance, parameters, healthcare, distribution

1. Introduction

The purpose of this study is to showcase how specific AI mechanisms can be introduced to transform raw medical insurance data into a strategic asset. An actuality over the world, not just in the United States, is that the population suffers from more afflictions with each passing day due to various economical, social and political changes. This global shift toward a higher disease burden necessitates a transition from reactive care to predictive financial modeling. The rise of medical insurance comes from an inability to pay healthcare services in a singular act. Insurance is the safest way to hold onto consistent healthcare services, as it is proven that an interruption in critical services leads to unfortunate consequences for patients and providers alike.

The healthcare landscape in the United States is among the most complex in the world, characterized by a diverse population with widely varying health profiles, income levels, and geographic access to care. According to the Centers for Disease Control and Prevention, chronic conditions such as obesity, diabetes, and cardiovascular disease continue to rise, placing enormous financial pressure on both insurers and policyholders. These structural pressures make accurate risk modeling not merely a competitive advantage but an operational necessity for insurance firms hoping to remain solvent over the long term.

While the current industry standard for AI implementation revolves heavily around customer satisfaction, claim processing, fraud detection and risk mitigation, this study argues that these are merely the minimum requirements for modern firms. To achieve a true competitive advantage, insurance companies must look beyond data organization and toward strategic optimization. By analyzing the connection between risk factors, the results compress the impact of each factor to the insurance charge paid by a person. The U.S. health insurance market benefits from a demographic complexity that is unmatched globally. This diversity provides a rich foundation for training machine learning models that can account for a wide spectrum of health profiles.

Scalability and equity are two components that define the implementation of AI in the insurance market, further increasing competitive advantage to companies that adapt faster. The data collected must therefore reflect the demand on the market for certain medical insurance plans that companies can categorize and upgrade in real time with the help of AI. This study focuses on multiple aspects that influence the demand for different insurance plans: smoker status implies a higher risk in terms of diseases of the lungs or heart, while Body Mass Index indicates targeted needs for people in the obesity range. Age, as a continuous variable, adds another layer of granularity that static actuarial tables fail to capture in a dynamic manner.

This rigorous quantitative approach serves to validate whether AI-driven predictions consistently outperform traditional linear actuarial methods in maintaining a firm's profitability and market position. By establishing a clear empirical baseline through statistical analysis, this paper creates a bridge between existing academic literature and the practical implementation strategies that insurance firms must consider as they modernize their operations. The conclusions drawn here are not speculative projections but evidence-based recommendations grounded in real demographic data from the United States.

In order to better understand how AI is being integrated in the health service business, relevant scientific papers from the last three years have been selected to contextualize the findings of this study. As for how this paper expands the view on AI implementation, based on diverse hypotheses, results have been analyzed through statistical tests on random samples drawn from a publicly available medical insurance dataset. The intersection of data science and insurance economics is where this paper positions itself, offering both methodological clarity and strategic insight.

2. Literature Review

Artificial Intelligence has already been incorporated and is being used at a high capacity by insurance companies. Health-related AI applications such as Hippocratic AI, Merative, Viz.ai, Enlitic, Regard, Twill, Linus Health, PathAI, VirtuSense, Cleerly, and Freenome are already specialized in different areas of medical services, from data analysis and image interpretation to valuable aid in neuroscience (Ige, 2024). While these algorithms are trained to assist in complex diagnostics, a more popular integration of AI involves chatbots and claim processing, yet rapid advances are being made toward cross-selling opportunities, marketing individualization, and the prediction of potential losses (Alkhelb and Alshagrawi, 2025).

Nellutla (2025) showcases how focusing on a single part of AI integration brings forth markedly different results depending on organizational context and implementation maturity. Case studies of Anthem and UnitedHealthcare demonstrate how AI drives competitive advantage in concrete, measurable ways. Anthem utilized Robotic Process Automation and Natural Language Processing to automate workflows, reducing processing times by 40% and eliminating approximately five million dollars in fraudulent claims. UnitedHealthcare employed predictive analytics to improve approval accuracy by 30%. Together, they illustrate that integrating operational automation with strategic data forecasting can simultaneously optimize efficiency, resource allocation, and customer satisfaction across an organization.

To maintain a sustainable competitive advantage, insurance firms must implement robust defensive mechanisms alongside their growth strategies. As outlined by Pandya (2023), the integration of machine learning and deep learning frameworks allows for the real-time detection and prevention of health insurance fraud. By automating the identification of anomalous claim patterns, insurers can significantly reduce losses that would otherwise erode profitability. This defensive layer of AI functionality is not optional in today's environment; it is foundational to protecting the revenue that more sophisticated AI applications seek to optimize.

The literature also highlights the complex relationship between age-based coverage policies and health outcomes more broadly. Yoruk and Han (2024) examine how age-based insurance coverage eligibility affects mental health outcomes, demonstrating that financial access to healthcare has measurable impacts beyond physical illness. This finding reinforces the importance of designing insurance plans that account not only for physical risk factors but also for the psychological and behavioral dimensions of health that bear on long-term costs. AI systems trained on multi-dimensional datasets that include behavioral and mental health indicators are therefore likely to outperform those trained on physical markers alone.

However, the pursuit of competitive advantage through AI is not without systemic risks. As described by Mello et al. (2026), the current AI arms race in the insurance sector, specifically within utilization reviews, presents a dual reality. On one hand, AI promises efficiency gains and more accurate coverage decisions; on the other, it risks codifying existing biases into automated systems that operate at enormous scale. This concern is not hypothetical. If a training dataset reflects historical disparities in care access or pricing, the resulting model will perpetuate those disparities with algorithmic authority. The studies reviewed therefore collectively point to a need for AI implementations that are not only powerful but transparent, auditable, and continuously monitored for fairness.

What the existing literature has not fully addressed is the precise quantitative relationship between demographic variables and insurance charges as a basis for AI model development. Most studies focus on a single application domain, such as fraud detection or claim automation, rather than examining the foundational statistical relationships that determine how well any AI system will perform when applied to pricing and risk stratification. This study seeks to fill that gap by establishing a clear statistical baseline that can inform future AI development efforts in the insurance sector.

3. Methodology

3.1 Data Collection

This study utilizes the Medical Insurance Cost Dataset (Abdelghany, 2024), representing a demographic cross-section of beneficiaries in the United States. The dataset consists of 1,338 individual records with features including Age, Sex, BMI, Number of Children, Smoking Status, and Geographic Region. This sample size, while not exhaustive, is sufficiently large to produce statistically meaningful results across all major demographic categories present in the U.S. insurance market.

Preliminary data cleaning was conducted to identify null values and handle potential duplicates, specifically investigating records with a BMI of 30.59 and charges of 1,639.56, to ensure the integrity of the statistical output. The deduplication process confirmed that no systematic data entry errors compromised the dataset's reliability. The Number of Children variable was set aside in the primary analysis of how each parameter functions as a significant determinant in AI implementation, since childbirth and dependent care represent a separate actuarial discussion requiring more specialized modeling frameworks. The remaining variables, age, sex, BMI, smoking status, and geographic region, were retained as the primary predictors of medical insurance charges.

The geographic variable encompasses four U.S. regions: northeast, northwest, southeast, and southwest. This regional categorization is meaningful because healthcare costs, provider availability, and the prevalence of chronic conditions vary systematically across these geographic zones.

3.2 Statistical Analysis

The analysis was performed using a Python-based framework leveraging pandas for data manipulation, seaborn and matplotlib for visualization, and scipy along with statsmodels for inferential testing. The methodology follows a tiered approach designed to move from descriptive understanding through diagnostic validation and into predictive modeling, each stage informing the next.

The first tier consists of Exploratory Data Analysis, in which all variables are visualized to identify distributional properties including skewness, kurtosis, and the presence of outliers. This phase establishes the empirical character of the data before any inferential claims are made. The second tier involves Assumption Testing through the Shapiro-Wilk test for normality and Levene's test for equality of variance. These tests are critical because they determine which inferential methods are appropriate given the data's actual distributional properties.

The third tier applies Hypothesis Testing through Independent Samples T-Tests, comparing groups such as smokers and non-smokers, as well as One-Way Analysis of Variance to examine whether regional differences in charges are statistically meaningful. The fourth tier involves Correlation and Regression Analysis, combining Pearson and Spearman coefficients to detect both linear and monotonic relationships, and ultimately constructing a Multiple Linear Regression model that quantifies the predictive weight of each demographic and behavioral factor. This layered approach reflects the rigor expected of a study intended to inform applied AI development decisions in a commercially sensitive domain.

4. Results and discussions

The analysis of the insurance dataset, conducted through multiple visual and statistical lenses, provides a comprehensive understanding of the factors driving medical costs and offers actionable insights for improving AI-driven insurance strategies. The initial exploration begins with seven graphs grouped together in order to establish a clear overview of all parameters simultaneously (Figure 1).

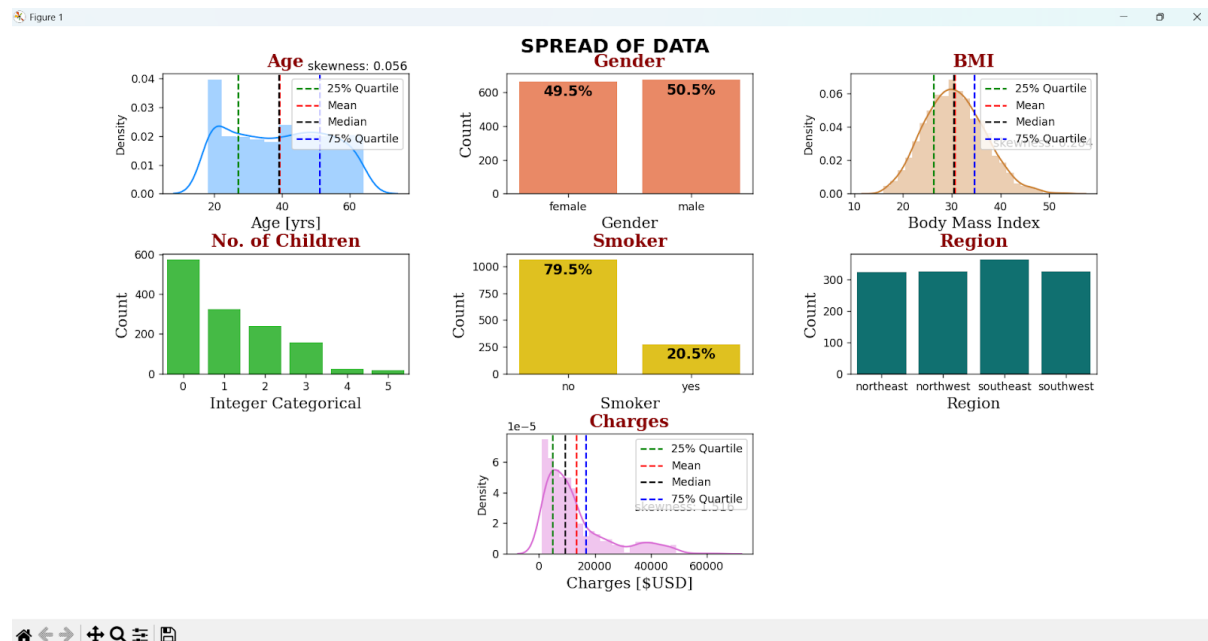


Figure 1

Source: Own configuration through code.

The age distribution reflects a broad demographic reach across the sample, spanning young adults through older individuals in a manner that mirrors the general structure of the insured U.S. population. The BMI data approximates a normal distribution with a slight right-hand skew, indicating that a significant portion of the population sits at higher health risk thresholds. This skew is consequential for insurance modeling because individuals with elevated BMI are more likely to present with conditions such as type 2 diabetes, hypertension, and joint disorders, all of which translate into higher claims over time.

The gender distribution is nearly even across the dataset, which is significant for ensuring a fair and unbiased analysis across the population. An imbalanced gender distribution could introduce systematic distortions into any model trained on the data, so the approximate parity observed here adds confidence to the generalizability of the results. Similarly, the regional distribution shows only a modest overrepresentation of the southeast region, a minor asymmetry that does not meaningfully compromise the integrity of cross-regional comparisons.

The most critical finding in this initial phase is the distribution of medical charges, which exhibits high positive skewness. This long tail in the charge distribution represents a small group of policyholders who incur disproportionately high costs, presenting a major challenge for traditional linear models that assume a roughly symmetrical distribution of outcomes. For an AI system to be effective in this environment, it must be capable of identifying the risk factors associated with these high-cost outliers so that health insurance companies can develop targeted products that remain profitable even when serving this high-risk segment. Companies that can do this accurately gain a structural advantage over competitors who rely on blunt, population-level pricing.

Moving into statistical diagnostics, the normality of insurance charges is rigorously tested using Q-Q plots and the Shapiro-Wilk test, as illustrated in Figure 2.

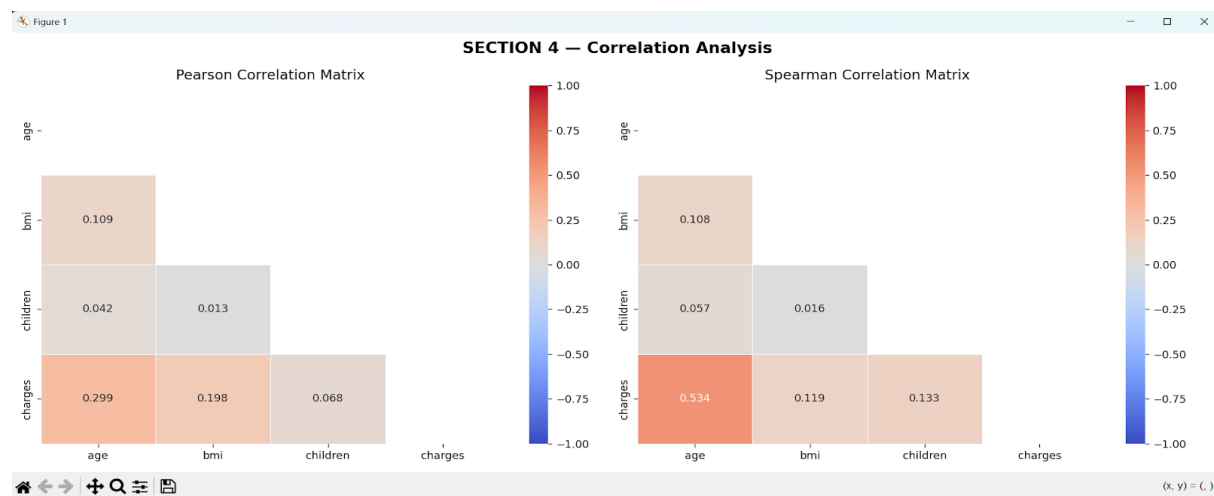


Figure 2

Source: Own configuration through code.

The Q-Q plot reveals significant deviations from the theoretical normal distribution, particularly in the upper quartiles, with a p-value of less than 0.05 confirming the departure from normality. This non-normality is not merely a statistical inconvenience; it carries direct implications for model selection. Simple linear assumptions are insufficient for capturing the complexity of healthcare spending across the full population, particularly in the high-cost tail where the most consequential pricing decisions must be made. Simultaneously, Levene's test for equality of variances demonstrates that the variance in charges between smokers and non-smokers is significantly different, a finding known as heteroscedasticity that further disqualifies ordinary least squares regression as a definitive modeling solution.

These diagnostic steps are essential for AI model selection, as they provide the empirical justification for moving toward non-linear algorithms, gradient-boosting machines, or ensemble methods that can better handle the heteroscedastic and skewed nature of the target variable. A practitioner who skips this diagnostic phase risks deploying a model that performs adequately on average but fails catastrophically in exactly the cases that matter most for profitability, namely the high-cost, high-risk individuals whose claims can determine whether a book of business is profitable or not.

The impact of specific risk factors is further examined through independent samples t-tests, most notably when comparing smokers to non-smokers. The resulting boxplots in Figure 3 show a dramatic disparity, with the mean charges for smokers far exceeding those of non-smokers across every demographic subgroup in the sample.

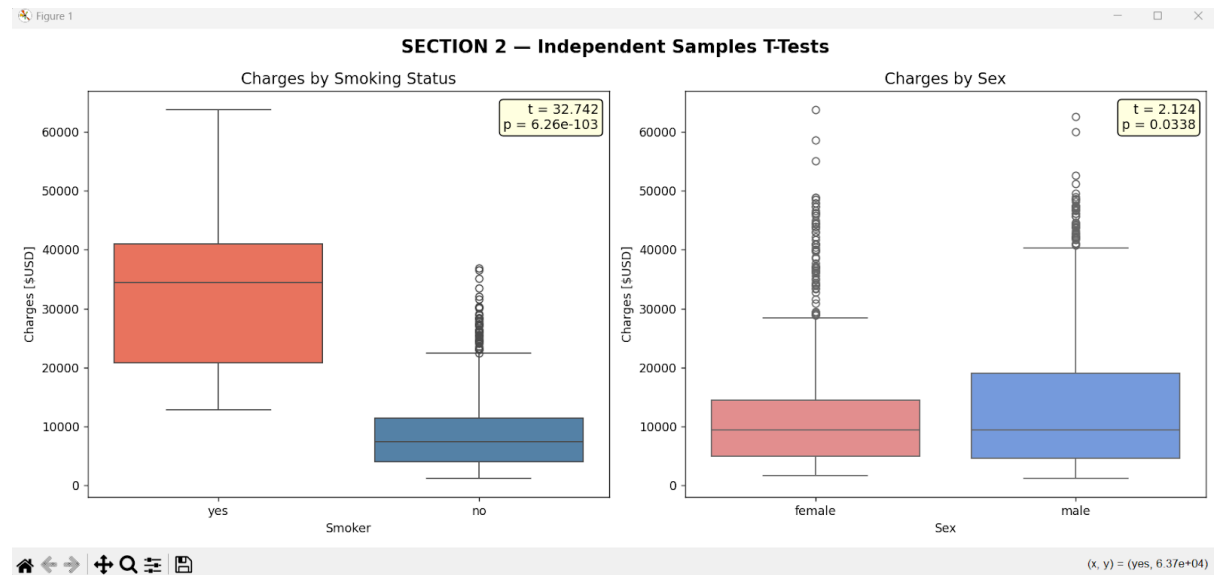


Figure 3

Source: Own configuration through code.

This result confirms smoking as a primary cost driver and establishes the empirical basis for more aggressive risk-weighting strategies targeting tobacco users. The median cost for smokers resides on an entirely different financial plane, often tripling or quadrupling the costs incurred by non-smokers with otherwise similar profiles. This clear separation is statistically validated by an exceptionally high t-statistic, confirming that the cost differential is not a result of random sampling variation but a systemic and reliable phenomenon that any competent AI model must detect and price accordingly.

In contrast, the comparison of charges by sex typically shows a much smaller, often non-significant difference. The interquartile ranges overlap substantially, and the whiskers of the boxplots extend to similar cost boundaries across both groups, indicating that gender does not carry the same categorical predictive weight as smoking status. While minor fluctuations in mean costs exist between male and female policyholders, they frequently fail the test of statistical significance in a predictive context.

One-Way ANOVA and Tukey HSD post-hoc tests further investigate regional differences, examining whether geography plays a meaningful and actionable role in determining medical costs across the four U.S. regions represented in the dataset (Figure 4).

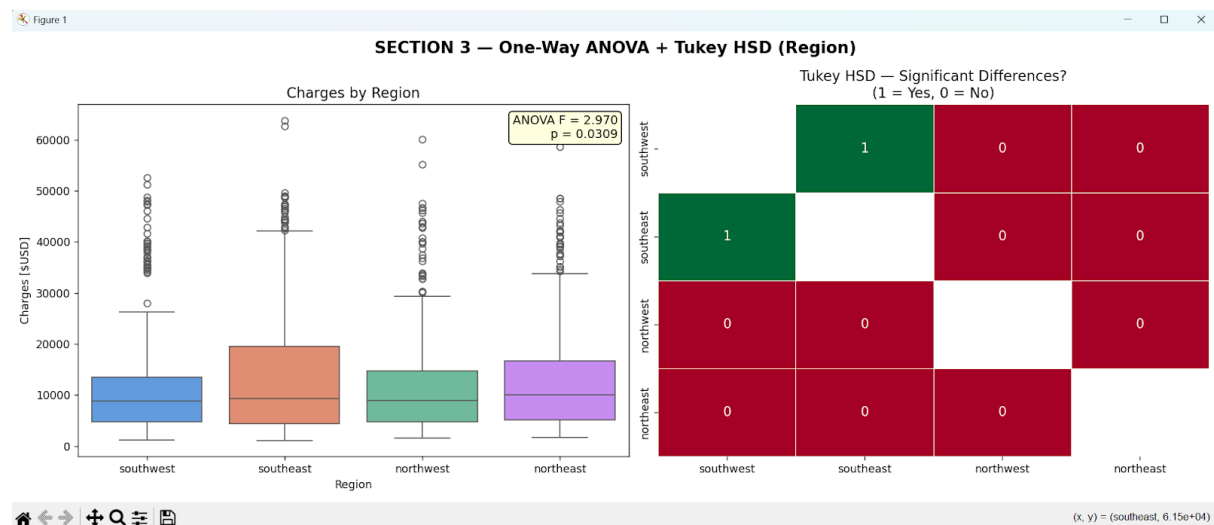


Figure 4

Source: Own configuration through code.

The spatial analysis of insurance data provides a detailed view of how geographical position influences financial liability in ways that categorical factors like smoking status do not reveal. While smoking produces massive and obvious shifts in cost, regional variations are subtle, requiring the precision of AI-driven analysis to determine whether they warrant a region-specific pricing strategy. The most informative visualization in this section is the Tukey HSD heatmap, which provides a better comparison of all regional groupings.

The heatmap reveals that while most regions share broadly similar cost profiles, certain pairings, most notably the Southeast versus the Northwest, show statistically significant mean differences. These disparities can often be traced to localized factors including variations in the cost of living, regional healthcare infrastructure quality, the density of specialist providers, and even the prevalence of chronic conditions associated with local dietary patterns or environmental exposures. For an AI system deployed at scale, these regional signals provide an additional layer of pricing granularity that, while smaller in magnitude than smoking status or BMI, can meaningfully improve the calibration of risk scores when combined with other variables.

Figure 5 presents the correlation matrices, using both Pearson and Spearman methods to quantify the relationships among the continuous variables: age, BMI, and charges.

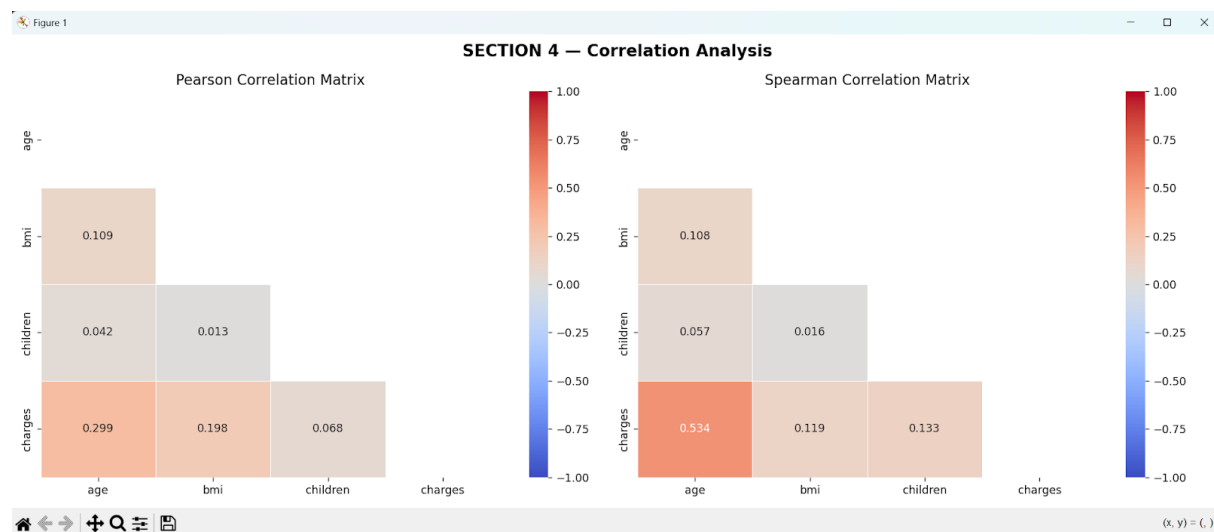


Figure 5

Source: Own configuration through code.

The Pearson heatmap identifies linear correlations, while the Spearman matrix captures monotonic relationships that may be non-linear in character. Age consistently shows a moderate positive correlation with charges, reflecting the natural increase in health risks over time as organ systems become less resilient and chronic conditions accumulate. This relationship is well understood in actuarial science, but the key question for AI implementation is not whether the relationship exists but how precisely it can be modeled across the full age spectrum.

BMI presents a more complex, non-linear challenge for traditional actuarial models. The Pearson heatmap reveals a positive linear correlation between BMI and charges, but the Spearman matrix uncovers a stronger monotonic relationship that suggests the association accelerates rather than grows at a constant rate. This is consistent with clinical evidence showing that as BMI crosses specific thresholds, particularly the transition into obesity at a BMI above 30, the associated medical costs do not merely increase incrementally but can accelerate sharply. Conditions like sleep apnea, metabolic syndrome, and osteoarthritis cluster disproportionately in the obese range, creating a non-linear cost signature that a linear model will be underestimated.

The interaction effects between these variables become especially visible when charge data is visualized through scatter plots colored by smoking status (Figure 6). Charges rise with age for everyone, yet the baseline and the rate of increase are significantly higher for smokers at every age point.

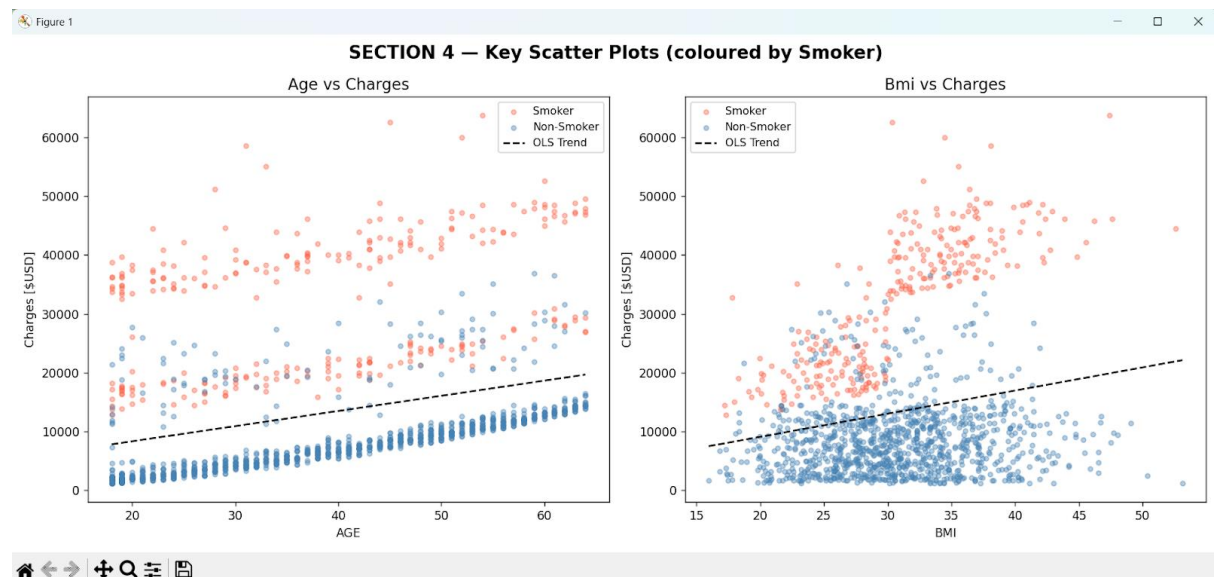


Figure 6

Source: Own configuration through code.

These interaction effects, where the impact of one variable depends on the value of another, are precisely what an advanced AI model must capture to provide accurate individual-level risk assessments. A model that treats age, BMI, and smoking status as independent additive factors will underprice the risk of smokers of old age with elevated BMI, leaving the insurer financially exposed on exactly the policyholders most likely to generate high claims.

By integrating these multi-dimensional relationships, insurance providers can shift from broad demographic pricing to highly personalized strategies that more accurately reflect an individual's actual risk profile. This shift is consequential not only for profitability but also for competitive positioning: a company that can offer competitively priced products to low-risk individuals while accurately pricing high-risk individuals will outperform competitors that apply uniform adjustments across risk categories.

Figure 7 presents three diagnostic and inferential graphs from the Multiple Linear Regression analysis, each serving a distinct analytical purpose.

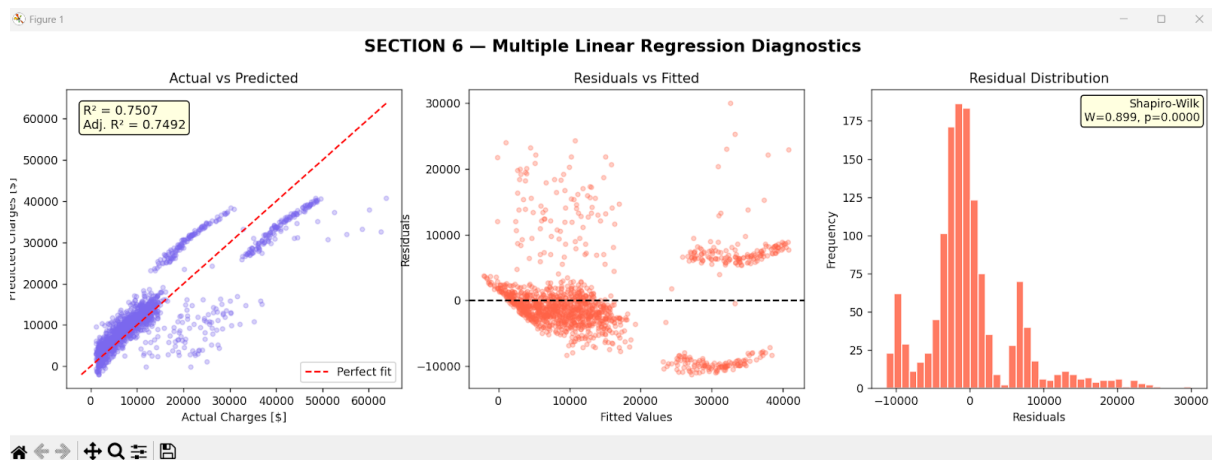


Figure 7

Source: Own configuration through code.

The Actual vs. Predicted scatter plot, together with the reported R-squared value of 0.7509, provides a direct measure of the model's explanatory power. The model successfully accounts for approximately 75% of the variance in medical costs using only six routinely collected demographic and behavioral variables. In a scientific context, this result is considered highly robust for a linear model applied to real-world healthcare data, which is inherently noisy and subject to unpredictable individual-level variation. The remaining 25% of unexplained variance likely stems from factors not captured in the dataset, including unpredictable catastrophic health events, rare genetic conditions, or the timing and quality of care accessed by individual policyholders.

The Residuals vs. Fitted plot serves as a diagnostic safeguard for the integrity of the regression model. In an ideal model, the residuals should be randomly dispersed around the horizontal zero-line, indicating that errors are attributable to random noise rather than systematic bias. Where the residuals show patterns, particularly the fan-shaped pattern associated with heteroscedasticity in the higher fitted values, this confirms the earlier diagnostic findings and reinforces the case for non-linear model architectures in production AI systems.

The Regression Coefficients graph functions as the strategic command center for insurance pricing optimization. The coefficient for smoking status carries a magnitude of nearly 23,848 dollars per individual with a tight confidence interval, meaning that the model's estimate of the smoking cost premium is both large and precise. This precision matters enormously in practice: a confident, well-bounded coefficient can be used directly to justify pricing decisions in regulatory filings, whereas a wide confidence interval would introduce unacceptable uncertainty into those decisions. By providing 95% confidence intervals for all predictors, the regression output gives the AI system a measure of pricing stability that static actuarial tables cannot match, because the confidence intervals are derived from the specific data distribution rather than population-level assumptions.

This data-driven pricing framework allows the firm to offer more competitive rates to individuals classified as low-risk, such as young non-smokers with healthy BMI values who represent stable, high-margin revenue streams, while ensuring that higher-risk individuals pay premiums that are mathematically calibrated to cover their expected medical expenditures. The ethical dimension of this approach is significant: low-risk individuals currently subsidize the healthcare costs of high-risk individuals in poorly differentiated pricing schemes. AI-enabled risk stratification allows insurers to construct more equitable portfolios where premiums more accurately reflect individual risk, creating a market incentive for healthier behaviors while reducing the cross-subsidization burden on low-risk policyholders.

For the insurer, identifying patterns in residual dispersion is vital because it highlights specific segments of the population, such as high-BMI smokers, where the linear model is systematically underestimating risk. By recognizing where the residual dispersion pattern is greatest, data scientists can implement non-linear adjustments or develop specialized sub-models to capture these extreme outliers, thereby protecting the overall portfolio from unexpected losses that a single linear model cannot adequately anticipate.

5. Conclusions

The findings of this study do not merely characterize the statistical determinants of healthcare costs; they make a compelling empirical case for why artificial intelligence must become a cornerstone of the American healthcare insurance system. The regression model developed here, achieving an R-squared of approximately 0.75 using only six routinely collected variables, demonstrates that medical expenditure is far from unpredictable. It is structured, quantifiable, and critically modelable with tools that are already available to any insurance company willing to invest in the necessary analytical infrastructure.

This predictability is the foundational premise upon which AI-driven healthcare transformation rests. If costs were truly random, no amount of data collection or algorithmic sophistication could improve pricing decisions. The fact that three-quarters of the variance in charges can be explained by basic demographic and behavioral variables confirms that the signal exists and is strong enough to support both commercially viable and ethically defensible AI applications in insurance pricing and product design.

The dominance of smoking status as a cost predictor, carrying a coefficient of nearly 23,848 dollars per individual with a remarkably tight confidence interval, illustrates precisely the kind of high-signal pattern that machine learning systems are designed to exploit at scale. Where a human actuary reviews individual cases sequentially, an AI risk-stratification engine can process millions of policyholders simultaneously, flagging high-risk individuals before costs are incurred rather than after. This proactive posture is fundamentally different from the reactive model that has historically defined insurance operations, and it represents a structural improvement in both efficiency and financial performance.

The age and BMI coefficients, though smaller in magnitude than smoking status, are equally actionable for strategic purposes. They define a continuous risk gradient that AI models can use to generate dynamic, individualized risk scores updated in real time as patient data evolves through connected health records, wearable devices, or periodic reassessments. This

dynamism is something no static actuarial table can achieve, and it represents the most important qualitative difference between traditional insurance pricing and AI-enabled pricing. A policy priced on the basis of a five-year-old health assessment may be significantly miscalibrated relative to the policyholder's current risk profile; an AI system with access to current data can recognize and correct for this drift continuously.

Equally significant is what the model reveals about the relative importance of gender and region as pricing factors. The near-zero, non-significant coefficients for sex confirm that the primary cost drivers in this dataset are behavioral and physiological rather than demographic in the traditional sense. A well-trained AI system built on this evidence base would therefore allocate risk and resources on clinically meaningful grounds rather than perpetuating the demographic generalizations that have historically created both pricing inaccuracies and legal exposure for insurance firms. This is not a peripheral benefit of AI adoption; it is central to the long-term regulatory and reputational sustainability of any insurance business that deploys these tools.

The approximately 25% of variance left unexplained by the linear model represents both a current limitation and a clear developmental opportunity. It is precisely this gap that advanced AI architectures are positioned to close. Gradient boosting models, deep neural networks, and ensemble methods capable of capturing the non-linear interactions between smoking, age, and BMI can explain a substantially higher proportion of the variance than the linear baseline established here. When these models are further enriched with electronic health record data, prescription histories, longitudinal biometric measurements from wearables, and behavioral data from lifestyle applications, they have the potential to approach near-complete cost predictability at the individual level for the large majority of policyholders. The present analysis establishes the statistical foundation. The logical and necessary next step for the industry is to build upon it systematically through AI.

The American healthcare system faces a structural cost crisis that no incremental administrative reform has proven capable of resolving. This study demonstrates, with statistical rigor, that the data required to address it already exists, that the predictive signal is strong enough to support meaningful model development, and that the methodological pathway from raw data to actionable AI is clear and reproducible. Artificial intelligence is not a speculative or experimental solution in this context. It is the most direct and evidence-grounded route from the current state of reactive, imprecise insurance pricing to a future healthcare system that is more efficient in its allocation of resources, more equitable in its treatment of policyholders, and more financially sustainable for the firms and consumers who depend on it.

6. References

- Abdelghany, M. (2024) Medical Insurance Cost Dataset. Available at: <https://www.kaggle.com/datasets/mosapabdelghany/medical-insurance-cost-dataset>
- Ige, T.O. (2024) 'An analysis of the use of Artificial Intelligence in the United States of America's healthcare system', *Journal of Health, Medicine and Nursing*, 113, pp. 59-67. doi: 10.7176/JHMN/113-07
- Nellutla, A. (2025) 'Enhancing Health Insurance Claim Processing through Artificial Intelligence and Data Analytics', *International Journal of Novel Research and Development*, 10(1), pp. c542-c550. Available at: <https://www.ijnrd.org/viewpaperforall?paper=IJNRD2501264> (Accessed: 9 March 2026).
- Alkhelb, A.A. and Alshagrawi, S. (2025) 'Role of artificial intelligence in healthcare insurance: systematic literature review', *Exploration of Digital Health Technologies*, 3, 101145. doi: 10.37349/edht.2025.101145.
- Pandya, S. (2023) 'A machine and deep learning framework for robust health insurance fraud detection and prevention', *International Journal of Advanced Research in Science, Communication and Technology*, 3(3), pp. 1332-1341. doi: 10.48175/IJAR SCT-14000T
- Mello, M.M., Trotsyuk, A.A., Mahamadou, A.J.D. and Char, D. (2026) 'The AI arms race in health insurance utilization review: promises of efficiency and risks of supercharged flaws', *Health Affairs*, 45(1). doi: 10.1377/hlthaff.2025.00897.
- Yoruk, B.K. and Han, Y. (2024) 'Age-based health insurance coverage policies and mental health', *Journal of Population Economics*, 37(2). doi: 10.1007/s00148-024-01015-w.