

Medical Insurance premium prediction using machine learning and applying concepts of MLOPS for deploying.

Pedarasi Sreyanwika, Murahari Shiva Jyothi
Department of information and Science technology
Presidency University, Bangalore
SREYANWIK.A.20211BCA0044@presidencyuniversity.in,
MURHARI.20211BCA0021@presidencyuniversity.in
Asst.professor vivek bhogle, Asst.professor jarnail singh

Abstract Predicting health problems is a major research point for actuarial scientists in the healthcare industry, as the more insurance companies try to use machine learning (ML) techniques to increase productivity and efficiency, the more they focus on identifying health conditions. Machine learning (ML) methods are still being researched and developed in the healthcare industry, so accurately predicting health insurance costs is an area of active research and development. ML in healthcare includes many techniques to help people by aiding rapid diagnosis and prediction of diseases, which in most cases doctors are unable to do. The fusion of technologies, especially digital health insurance, will be the one to close the distance between insurers and consumers, they are directly connected. Machine learning is the reason why the conventional insurance system has been completely changed and this has led to the rapid delivery of services in creating policies that insurers have noticed and adopted. In this way, insurance companies can provide their clients with health insurance services that are fair, fast and efficient through ML. This study used and analyzed multiple ML-based regression models to determine health insurance premiums. The researchers calculated the predicted insurance costs for people based on their characteristics. They set standards for numerous factors like age, gender, BMI, number of children, etc. and in the process even set performance baselines for the same. Models were given and tested using these parameters. Among other experimental results, the xgboost regression model won with 90% accuracy. 22% of actions were effectively transferred from real life to simulation. The authors reported these results and evaluated model performance using key performance

metrics. The only thing is, having a model in a Jupyter notebook is just the beginning. Establishing a pipeline and review process through MLOps best practices is an emerging trend for many companies. The study concludes by demonstrating that the entire process of containerizing a Flask-based ML web application to its deployment on the Heroku cloud platform has been successfully completed.

Keywords

Machine learning, medical charges, prediction, Mlops, Flask, Procfile, Heroku.

INTRODUCTION

The world we live in is full of dangers and unpredictable factors. People, families, businesses and structures are the ones who have to face many hazards that can bring them different types of risks such as death, disease and loss of property or assets. Although people put their health and well-being first, it is not possible to completely eliminate risks. Numerous products that have been created by the financial sector to protect financial resources in the event of losses are measures that are taken to prevent the risks of financial crimes. Thus, insurance is a system put in place to reduce or completely eliminate expenses related to various types of risks. In particular, health insurance is a policy that covers medical expenses. Purchasing health insurance and paying a predetermined premium are the first steps to getting coverage. The cost of health insurance is affected by several factors, therefore there is always a difference in the premium for one person to another. For example, age plays a significant role: the issue of health

problems is more pronounced in older individuals than in younger ones. The costs of caring for the elderly are therefore higher and, as a result, the insurance premiums for them are higher than for the young. Because of the various factors that affect health insurance premiums, these costs vary from person to person. ML algorithms have been found to be effective in predicting high costs and most important patient expenditures. Therefore, insurance companies are now increasingly turning to ML to update their policies and premium settings. MLOps, or Machine Learning Operations, is the name of the routine activities that define the processes of deploying machine learning models into production, as well as their maintenance and monitoring. MLOps refers to all processes from the data pipeline to the production of machine learning models. However, some of the implementations are just deploying machine learning models, many companies on the other hand use MLOps throughout the ML development lifecycle. This includes steps such as exploratory data analysis (EDA), data pre-processing and model training. MLOps is a process that comes with the idea of automation, scalability, reproducibility, monitoring and management for a machine learning project at all stages.

BACKGROUND

A key element of the medical industry is health insurance. In addition, it is not easy to estimate health expenditures because most of the money comes from patients. A number of ML algorithms and deep learning techniques are used to predict the data. Training time and factor accuracy are assessed. The sum of the machine learning algorithms requires only a short training time. However, the prediction results from these techniques are not very reliable. Deep learning models can also detect hidden patterns, but their real-time use is limited by training time. The different regression models used in this report are linear regression, XGBoost regression, lasso regression, random forest regression, and gradient regression. The XGBoost regression was found to be the most accurate, accounting for a maximum accuracy of 90.2 percent. The main purpose of this research is to present a new method of estimating insurance costs and using a cloud-based web application where N number of people can access it at a given time to know their health insurance premiums.

LITERATURE REVIEW

1. In the field of healthcare, it is necessary to research and improve methods of estimating health insurance premiums using machine learning algorithms. The presented work deals with the computational intelligence method for predicting health insurance costs using various ML techniques. One of the many essays began with an analysis of the possible effects of using predictive algorithms on insurance pricing. Can this lead to a threat to the principle of mutual risk, and thus to the emergence of new forms of bias and insurance exclusions? In the second phase, the authors analyzed the change in the situation between the company and the insured when customers learned that the company was always updating and collecting information about their actual behavior.

2. The aim of the research conducted by van den Broek-Altenburg and Atherly was to find out the opinions of customers about health insurance by examining their tweets. The goal was to use sentiment analysis to find out how people perceive health insurance and health care providers. The authors used an API to collect Twitter posts with the words “health insurance” or “health plan” during the 2016–2017 US health insurance enrollment period. Insurance is a contract whose aim is to reduce or completely eliminate expenses associated with various dangers. The price of insurance is determined by a number of factors, which in turn are related to the creation of insurance plans. Machine learning can help the insurance industry in the process of increasing the efficiency of the formulation of insurance contracts.

3. Nidhi Bhardwaj and Rishabh Anand's work dealt with the issue of premium prediction based on human health data. Many algorithms have been reviewed and evaluated using regression analysis. The dataset was the one used to train the models, and the training output was the one used to make the predictions. Finally, the models were tested and their accuracy verified by checking that the predicted values match the actual data. The reliability of the models was verified and the results showed that multiple linear regression and gradient boosting algorithms were better than linear regression and decision trees. Gradient boosting was the method that was chosen as optimal in this case because it performed the same thing as

multiple linear regression but needed a fraction of the time.

4. Risk ratings in the life insurance industry are becoming the tools that define life insurance beneficiaries. Organizations apply screening techniques to reapply decisions and to produce risk-adjusted pricing insurance products. Screening could be done by creating a computer interface for faster processing by applications and/or services. This can be achieved through growing data and creating business intelligence. This study aimed to identify techniques for using predictive analytics to improve risk assessment of insurance companies based on the lives of their customers/members. The data analysis investigation was performed on a real data set consisting of almost a hundred features (data anonymization). The element selection process was carried out using a dimension reduction method to those elements that will ultimately allow the models to provide more accurate results.

5. Health insurance premiums result from private statistical methods and the complex models employed by insurance companies, which are not disclosed to the public. The main objective of this research is to determine if ML techniques can foresee the health insurance premium prices for the next year as well as from the contract parameters and business characteristics. The main intention of this paper is to use a reliable ML model to project the future patient treatment costs which will be based on the certain parameters. The outcomes of the simulation were used as a basis to find the factors that are the sources of the variation in the health expenditure of individuals.

6. The Japanese government has ordered insurance companies to develop a population health management strategy. A major part of this strategy assessment is cost estimation. A standard linear model is not suitable for prediction because an insured patient may have multiple conditions. Using a quantitative machine learning technique, we developed a healthcare cost forecasting model. The research looked at the stability of health care spending in a large state Medicaid program. Predictive ML algorithms used for cost projections with a particular focus on HCHN patients. The results of Yang et al. showed strong temporal relationships and the practical potential of machine learning for forecasting health care expenditures.

HCHN patients showed closer temporal relationships, making their birth predictions more reliable. The study attempted to review and add additional historical periods, resulting in improved predictions.

7. Many research papers have been published, such as the one written by Fursin, the author proposes a multifaceted view of the CK platform from MLOps to open APIs. ThoughC[CN] has remained in the proof-of-concept phase of development and is going through several upgrade phases. The next step will look for ways to complete this task, such as the development of JSON Object and API Meta reports for all CK plugins and frameworks. Meanwhile, the initial goal was to accelerate energy transactions based on criteria that were measurable, including latency, accuracy, speed, and reliability, which is evident in the case of Moreschini et al.

8. Renggli et al. introduced the MLOps framework for AI-intensive software systems that are constantly evolving. The purpose is to develop a self-service component driven by artificial intelligence that can be updated to improve the software through an upgrade. Their main goal is to better understand how MLOps work in commercial applications.

9. Several community members strongly recommend that MLOps include a data quality process as part of operations. However, the main purpose of these visualizations is to indicate overlapping data quality attributes that can be found in all stages of the ML process. Although they mention some limitations of the models, two upcoming methods are introduced that can be defined as the struggle and the role of the teacher, and in this discipline, ML explorations are becoming an inevitable habit. A research team consisting of Granlund et al. offered a realworld setup in two scenarios to gain insight into the data and model management lifecycle (MLOps). They discussed the challenges of scaling and integrating ML both in academia and in the field of use of the technology. The company also provides a unified strategic operation to support other stakeholders. Mäkinen et al. attempted to determine a benchmark that could demonstrate the necessity of adopting MLOps for current ML systems, and the extent to which software engineers still encounter and report

problems with numbers without considering the implementation part. In the current setup, we expect to integrate MLOps in the future.

Methodology

1. Dataset Description

Data for model cost prediction using the Kaggle platform was obtained by our team. The dataset contains two distinct categories: training data and test data are the sources of the dataset, each consisting of seven attributes as shown in Table one. Most of the allocated data resources are for testing, with approximately 20% dedicated to training. The training set is used as input for the regressor to develop a model for calculating annual health insurance costs, and the test set is the medium for testing and evaluating the performance of the regression model. For better understanding, the following table row is available with a comprehensive analysis of the data set.

Table 1: Overview of the Dataset

Attribute	Data Description
Age	The age of individual person
Sex	Sex of the person (Male, Female)
BMI	This is Body Mass Index
Children	Total number of children of the person have
Smoker	Whether the person is a smoker or not
Region	Where the person lives. Considering four regions (Southwest, Southeast, Northeast, Northwest)

There were 1338 rows and 7 columns in our data set. The charges variable, which has a float value, is our aim. Maximum number of individuals in our dataset range in age from 18 to 22.5, and the majority of them are male. Few have more than three children, and the majority of them have a BMI between 29.26 and 31.16. In this dataset, four main regions are taken into account: northeast, northwest, southeast, and southwest. The largest concentration of smokers is in the southeast, where 1064 out of 1338 people smoke. We'll investigate our information to determine how the

various factors are related. Our target column in this instance is "charges," which is dependent upon every other column. We shall first examine our dataset's statistical metrics

Table 2: statistical measurement

	age	bmi	children	charges
count	1338.000000	1338.000000	1338.000000	1338.000000
mean	39.207025	30.663397	1.094918	13270.422265
std	14.049960	6.098187	1.205493	12110.011237
min	18.000000	15.960000	0.000000	1121.873900
25%	27.000000	26.296250	0.000000	4740.287150
50%	39.000000	30.400000	1.000000	9382.033000
75%	51.000000	34.693750	2.000000	16639.912515
max	64.000000	53.130000	5.000000	63770.428010

2. Data Analysis:

There were 1338 rows and 7 columns in our data set. The charges variable, which has a float value, is our aim. Here are some data visualizations.

Pie chart for each categorical variable columns named 'sex', 'smoker', and 'region'.

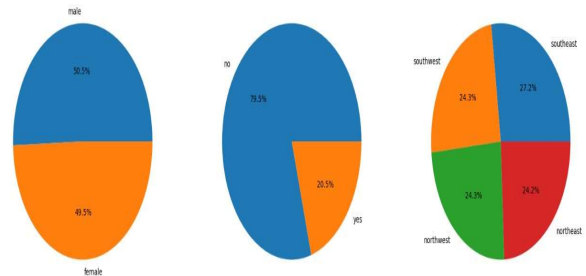


Figure 1. Pie chart about sex, smoker, region.

Figure 1 describes about an analysis of tweets from men in the Southeast showed a near 50:50 ratio of positive to negative sentiment, so the male population in the area appears to be neutral. 5% of people say yes and 49% of them say yes. 5% express negativity. Unlike their female counterparts in the Northeast region, men in the region express an almost exclusively positive attitude, represented by 48.5% expressing positivity, 27. ties with 2% negativity from the Southwest and 24, which is a highly connected region. 3% from the northwest. On the other hand, women in both places mostly show a negative mood, 79.3% of students are negative in Southeast and 51% of Southeast are involved in negativity. 5% of the Northeast are negative, but a smaller percentage show positive.

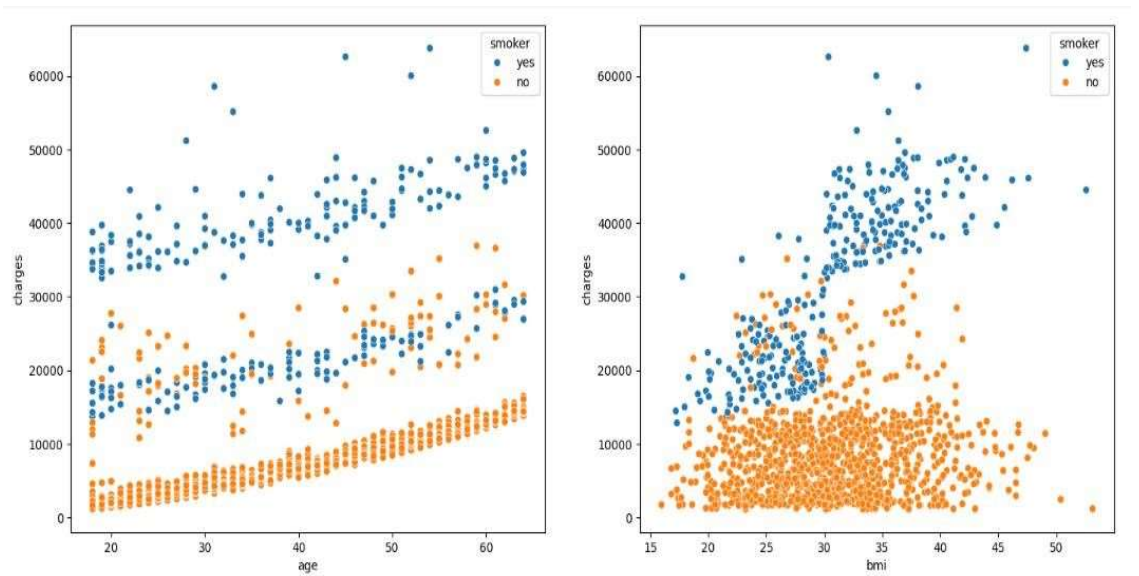


Figure 2. scatter plot of age, bmi

The figure shows two scatter plots showing the association between age and BMI. These two plots have a slightly upward trend. The left graph shows BMI on the y-axis and age on the x-axis, with ranges from 15 to 50 for BMI and 20 to 60 for age. A proper chart keeps the same axes but with a different color scheme. Text annotations have a "smoker" at the top left with "yes" and "no" options and numerical scales along the axes.

Below Line plot showing the relationship between 'charges' and 'bmi', as well as between 'charges' and 'age'.

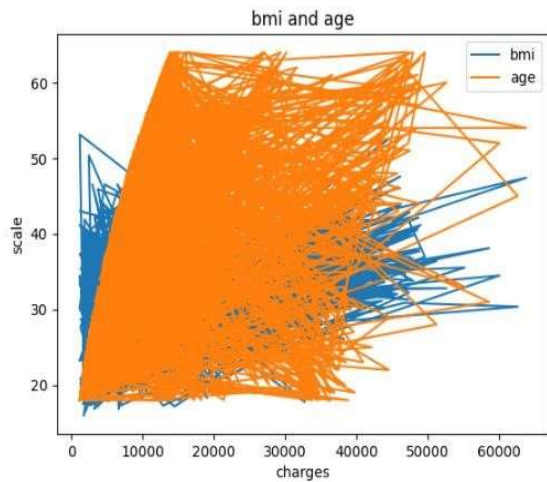


Figure 3: plotting relationship between charges and bmi as well as between charges and age.

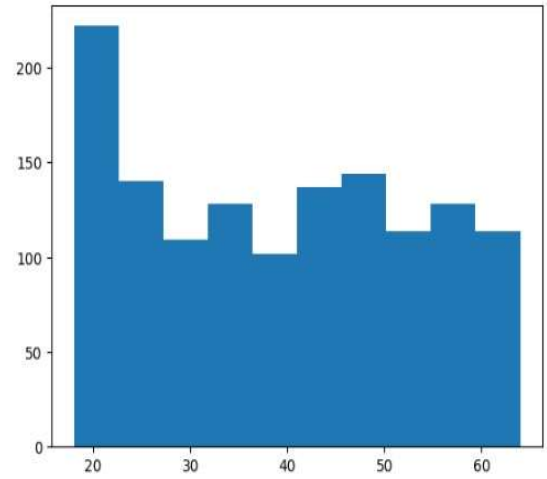


Figure 4: Hit plot for age

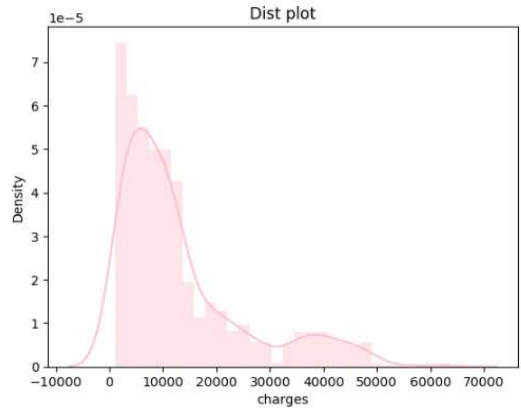


Figure 5: density plot for charges

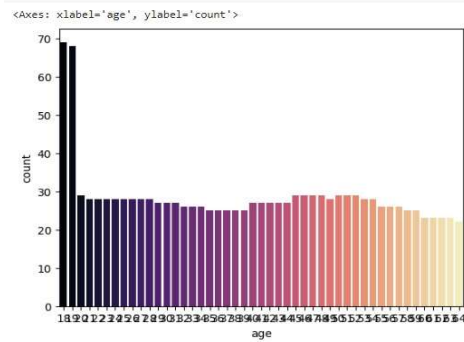


Figure 6: age distribution

Scatter plot with a linear regression line fit to the data points.

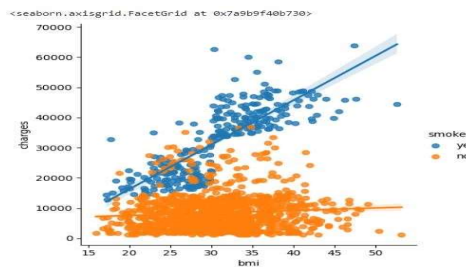


Figure 7: scatter plot with a linear regression line

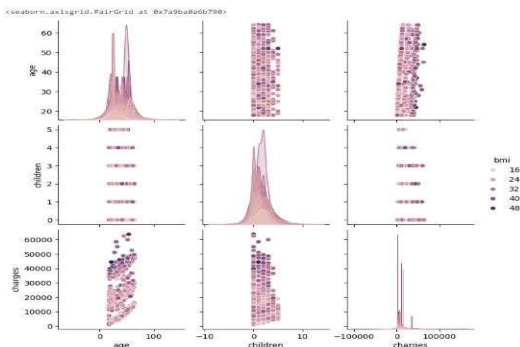


Figure 8: Visualization the relationship between two variables

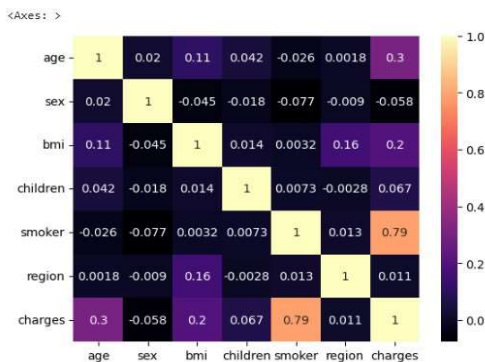


Figure 9: Heat map

3. Data Pre-processing

Three are columns that are numeric, and the other three are columns that are categorical. Our machine learning model cannot fit category values in text because computers cannot understand this text value. Thus, the categories are given a numerical designation as a result of the numerical assignment of the qualities. In the "gender" field, we change "female" to 1 and "male" to 0. In addition, we change the other two columns to be numeric. Our results are listed for conversion in the table below. **Table 3: Categorical to Numerical Conversion**

Column Name	Before Conversion	After Conversion
sex	male	0
	female	1
smoker	yes	0
	no	1
region	southeast	0
	southwest	1
	northeast	2
	northwest	3

4. Model specification

Table 4: Model Performance

Regression Models	R squared	RSME	MAE	MSE
Linear Regression	0.8062	5966.97	4190.09	35604738.22
Lasso Regression	0.8061	5967.76	4190.8	35614168.87
Random Forest Regression	0.882	4652.55	2586.47	21646283.48
Gradient Boosting regression	0.901	4249.84	2480.04	18061204.99
XGBoost Regression	0.9022	4270.73	2493.13	18239158.09

Our data is first obtained via Kaggle. Our dataset is then imported into Google Colab. Next, using various visualization tools, we analyze our data. The data is then cleaned such that it exactly matches the machine learning model. We then use our training data to apply regression techniques. Our model will be ready for cost forecasting after the data has been tested. The flowchart that follows illustrates the entire process.

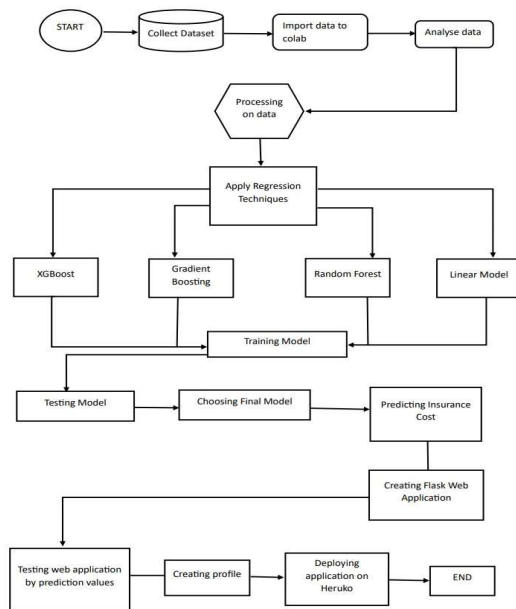


Figure-10: Flow Chart of Medical Insurance Cost Prediction System

Table-4 displays our top and bottom regression models. We can anticipate insurance costs using the model that performs best, according to the findings. In our situation, XGBoost Regression is the best regression model while Lasso Regression is the worst. Anyone may calculate their insurance expenses using the best model.

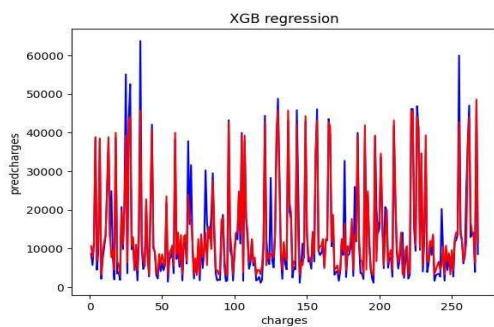


Figure 11: Predicted Cost using XGBBoost Regression

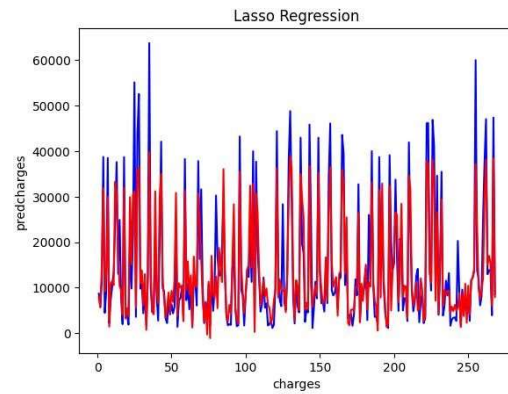


Figure 12: Predicted Cost Using Lasso Regression.

```

new_data=pd.DataFrame({'age':19,'sex':'male','bmi':27.9,'children':0,'smoker':'yes','region':'northeast'},index=[0])
new_data['smoker']=new_data['smoker'].map({'yes':1,'no':0})
new_data=new_data.drop(new_data[['sex','region']],axis=1)
finalmodel.predict(new_data)

array([[18835.828]], dtype=float32)
  
```

Figure 13: Medical Insurance cost prediction system demo result

5. Flask framework

The main part of the application is a Python file named app.py. It is built using the Flask Framework.

6. API (Application programming interface)

Development in Postman

Postman is an application programming interface (API) development tool that allows you to create, test, and modify APIs. This tool includes almost all the features that any developer can use. Postman software is a tool used to send and receive POST requests and data to a server.

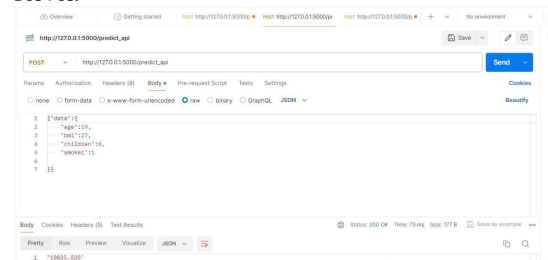


Figure 15: Results through postman

7. Procfile

A Procfile is a file used in some platforms, particularly in web development, to specify the commands that are executed by the application's dynos (containers) on the platform. It's commonly used in platforms like Heroku.



Figure 16: Procfile

8. Deploying on heroku

Heroku is a cloud platform as a service (PaaS) that enables developers to build, run and manage applications entirely in the cloud. It is backed by several programming languages like Ruby, Node, JS, Python, Java, PHP and so on. Heroku abstracts from infrastructure management, so developers can focus on building and deploying applications without thinking about hardware or server maintenance. It is widely used because it is easy to use, scalable and can be integrated with various development tools and services.

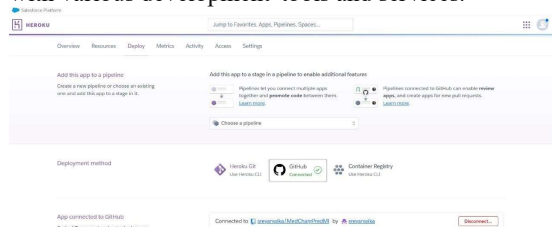


Figure 17: Connecting GitHub with Heroku to deploy application

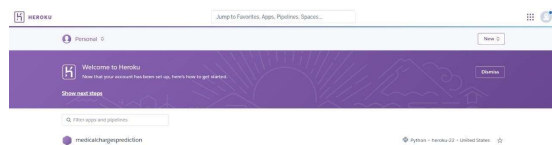


Figure 18: Deployed application

9. Results

We can anticipate insurance costs using the web application that performs best, according to the findings. In our situation, XGBoost Regression is the best regression model.

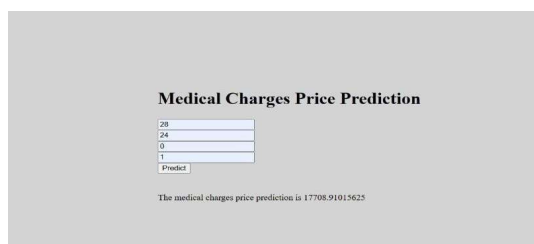


Figure 19: predicting charges through website.

Conclusion

In order to predict health insurance prices on Kaggle's individual health cost data set, the study uses ML

regression models. Table IV is a list of results. By calculating insurance rates based on multiple factors, insurance companies can attract customers and save time. Machine learning could substantially reduce this individual effort in cost analysis, as ML models can calculate costs very quickly, while it would take a human a long time. Large data can also be worked with using machine learning methods. Designing a finite model (xgb regression) web application using the flask framework and deploying it to the cloud using MLOps concepts makes it easy for every health insurance user and stakeholder. Mlops allows N users to access a website at once and also provides data security as user data is not stored or captured by the website. The work can be improved in the future by using a larger data set than the one used in this study, as well as better algorithms that perform better than XGB regression.

References

1. Cevolini A., Esposito E. From Pool to Profile: Social Consequences of Algorithmic Prediction in Insurance. *Big Data Soc.* 2020; **7** doi: 10.1177/2053951720939228. [CrossRef] [Google Scholar]
2. van den Broek-Altenburg E.M., Atherly A.J. Using Social Media to Identify Consumers' Sentiments towards Attributes of Health Insurance during Enrolment Season. *Appl. Sci.* 2019; **9**:2035. Doi: 10.3390/app9102035. [CrossRef] [Google Scholar]
3. Hanafy M., Mahmoud O.M.A. Predict Health Insurance Cost by Using Machine Learning and DNN Regression Models. *Int. J. Innov. Technol. Explor. Eng.* 2021; **10**:137–143. Doi: 10.35940/ijitee.C8364.0110321. [CrossRef] [Google Scholar]
4. Bhardwaj N., Anand R. Health Insurance Amount Prediction. *Int. J. Eng. Res.* 2020; **9**:1008–1011. doi: 10.17577/IJERTV9IS050700. [CrossRef] [Google Scholar]
5. Boodhun N., Jayabalan M. Risk Prediction in Life Insurance Industry Using Supervised Learning Algorithms. *Complex In tell. Syst.* 2018; **4**:145–154. Doi: 10.1007/s40747-018-0072- [CrossRef] [Google Scholar]
6. Goundar S., Prakash S., Sadal P., Bhardwaj A. Health Insurance Claim Prediction Using Artificial Neural Networks. *Int. J. Syst. Dyn. Appl.* 2020; **9**:Doi:10.4018/IJSDA.2020070103. [CrossRef] [Google Scholar]

7. Ejayi C.J., Qin Z., Salako A.A., Happy M.N., Nneji G.U., Ukwuoma C.C., Chikwendu I.A., Gen J. Comparative Analysis of Building Insurance Prediction Using Some Machine Learning Algorithms. *Int. J. Interact. Multimed. Artif. Intell.* 2022; 7:75–85. doi: 10.9781/ijimai.2022.02.005. [CrossRef] [Google Scholar]
8. Rustam Z., Yaurita F. Insolvency Prediction in Insurance Companies Using Support Vector Machines and Fuzzy Kernel C-Means. *J. Phys. Conf. Ser.* 2018; 1028:012118. doi: 10.1088/1742-6596/1028/1/012118. [CrossRef] [Google Scholar]
9. Fauzan M.A., Murfi H. The Accuracy of XGBoost for Insurance Claim Prediction. [(accessed on 9 May 2022)]; *Int. J. Adv. Soft Comput. Appl.* 2018 10:159–171. Available online: <https://www.claimsjournal.com/news/national/2013/11/21/240353.htm> [Google Scholar]
10. Kumar Sharma D., Sharma A. Prediction of Health Insurance Emergency Using Multiple Linear Regression Technique. *Eur. J. Mol. Clin. Med.* 2020;7:98–105. [Google Scholar]
11. Azzone M., Barucci E., Giuffra Moncayo G., Marazzina D. A Machine Learning Model for Lapse Prediction in Life Insurance Contracts. *Expert Syst. Appl.* 2022; 191:116261. doi: 10.1016/j.eswa.2021.116261. [CrossRef] [Google Scholar]
12. Sun J.J. Master's Thesis. National College of Ireland; Dublin, Ireland: 2020. [(Accessed on 9 May 2022)]. Identification and Prediction of Factors Impact America Health Insurance Premium. Available online: <http://norma.ncirl.ie/4373/> [Google Scholar]
13. Luis E. Employer Health Insurance Premium Prediction. [(Accessed on 17 May 2022)]. Available online: <http://cs229.stanford.edu/proj2012/LuiEmployerHealthInsurancePremiumPrediction.pdf>
14. Prediction of Health Expense—Predict Health Expense Data. [(Accessed on 9 May 2022)]. Available online: <https://www.analyticsvidhya.com/blog/2021/05/prediction-of-healthexpense/>
15. Takeshima T., Keino S., Aoki R., Matsui T., Iwasaki K. Development of Medical Cost Prediction Model Based on Statistical Machine Learning Using Health Insurance Claims Data. *Value Health.* 2018; 21:S97. doi: 10.1016/j.jval.2018.07.738. [CrossRef] [Google Scholar]
16. Yang C., Delcher C., Shenkman E., Ranka S. Machine Learning Approaches for Predicting High Cost High Need Patient Expenditures in Health Care. *Biomed. Eng. Online.* 2018; 17:131. doi: 10.1186/s12938-018-0568-3. [PMC free article] [PubMed] [CrossRef] [Google Scholar]
17. M. Chen, S. Mao, and Y. Liu, “Big data: A survey,” in *Mobile Networks and Applications*, Jan. 2014, vol. 19, no. 2, pp. 171–209. doi: 10.1007/s11036-013-0489-0.
18. Y. Lecun, Y. Bengio, and G. Hinton, “Deep learning,” *Nature*, vol. 521, no. 7553. Nature Publishing Group, pp. 436–444, May 27, 2015. doi: 10.1038/nature14539.
19. D. Baylor et al., “TFX: A TensorFlow-based production-scale machine learning platform,” in *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, Aug. 2017, vol. Part F1296, pp. 1387–1395. doi: 10.1145/3097983.3098021.
20. P. Agrawal et al., “Data platform for machine learning,” in *Proceedings of the ACM SIGMOD International Conference on Management of Data*, 2019, pp. 1803–1816. doi: 10.1145/3299869.3314050.
21. J. Deng, W. Dong, R. Socher, L.-J. Li, Kai Li, and Li Fei-Fei, “ImageNet: A large-scale hierarchical image database,” in *IEEE Conference on Computer Vision and Pattern Recognition*, Mar. 2009, pp. 248–255. doi: 10.1109/cvpr.2009.5206848.
22. E. Bisong, *Building Machine Learning and Deep Learning Models on Google Cloud Platform*. Apress, 2019. doi: 10.1007/978-1-4842-4470-8.
23. W. Hummer et al., “ModelOps: Cloud-based lifecycle management for reliable and trusted AI,” in *Proceedings - 2019 IEEE International Conference on Cloud Engineering, IC2E 2019*, Jun. 2019, pp. 113–120. doi: 10.1109/IC2E.2019.00025.
24. D. Sculley et al., “Hidden technical debt in machine learning systems,” in *Advances in Neural Information Processing Systems*, 2015, vol. 2015-Janua, pp. 2503–2511. Accessed: Jun. 08, 2022.

- [Online]. Available: <https://proceedings.neurips.cc/paper/2015/hash/86df7dcfd896fcdf2674f757a2463ebaAbstract.html>
25. B. Fitzgerald and K. J. Stol, "Continuous software engineering and beyond: Trends and challenges," in 1st International Workshop on Rapid Continuous Software Engineering, RCoSE 2014 - Proceedings, Jun. 2014, pp. 1–9. doi: 10.1145/2593812.2593813.
 26. B. Fitzgerald and K. J. Stol, "Continuous software engineering: A roadmap and agenda," *Journal of Systems and Software*, vol. 123, pp. 176–189, Jan. 2017, doi: 10.1016/j.jss.2015.06.063.
 27. L. E. Lwakatare, P. Kuvaja, and M. Oivo, "Dimensions of devOps," in *Lecture Notes in Business Information Processing*, 2015, vol. 212, pp. 212–217. doi: 10.1007/978-3-319-186122_19.
 28. T. Mikkonen, J. K. Nurminen, M. Raatikainen, I. Fronza, N. Mäkitalo, and T. Männistö, "Is Machine Learning Software Just Software: A Maintainability View," in *Lecture Notes in Business Information Processing*, 2021, vol. 404, pp. 94–105. doi: 10.1007/978-3-030-658540_8.
 29. L. Leite, C. Rocha, F. Kon, D. Milojicic, and P. Meirelles, "A survey of DevOps concepts and challenges," *ACM Computing Surveys*, vol. 52, no. 6. ACM PUB27 New York, NY, USA, Nov. 14, 2019. doi: 10.1145/3359981.
 30. D. A. Meedeniya, I. D. Rubasinghe, and I. Perera, "Software artefacts consistency management towards continuous integration: A roadmap," *International Journal of Advanced Computer Science and Applications*, vol. 10, no. 4, pp. 100–110, 2019, doi: 10.14569/ijacsa.2019.0100411.
 31. D. A. Meedeniya, I. D. Rubasinghe, and I. Perera, "Traceability establishment and visualization of Software Artefacts in DevOps Practice: A survey," *International Journal of Advanced Computer Science and Applications*, vol. 10, no. 7, pp. 66–76, 2019, doi: 10.14569/ijacsa.2019.0100711.
 32. I. Rubasinghe, D. Meedeniya, and I. Perera, "SAT-Analyser Traceability Management Tool Support for DevOps," *Journal of Information Processing Systems*, vol. 17, no. 5, pp. 972–988, 2021, doi: 10.3745/JIPS.04.0225.
 33. M. Virmani, "Understanding DevOps & bridging the gap from continuous integration to continuous delivery," in 5th International Conference on Innovative Computing Technology, INTECH 2015, Jul. 2015, pp. 78–82. doi: 10.1109/INTECH.2015.7173368.
 34. I. Rubasinghe, D. Meedeniya, and I. Perera, "Automated InterArtefact Traceability Establishment for DevOps Practice," in *Proceedings - 17th IEEE/ACIS International Conference on Computer and Information Science, ICIS 2018*, Sep. 2018, pp. 211–216. doi: 10.1109/ICIS.2018.8466414.
 35. D. Meedeniya and H. Thennakoon, "Impact Factors and Best Practices to Improve Effort Estimation Strategies and Practices in DevOps," in *ACM International Conference Proceeding Series*, Aug. 2021, pp. 11–17. Doi: 10.1145/3484399.3484401.
 36. A. Paleyes, R.-G. Urma, and N. D. Lawrence, "Challenges in Deploying Machine Learning: a Survey of Case Studies," *ACM Computing Surveys*, Nov. 2022, doi: 10.1145/3533378.