

# The Use of Artificial Intelligence in Gravitational Microlensing Detection for Dark Matter Discoveries

Suhaan Khan<sup>1\*</sup>

1. McNeil High School, 5720 McNeil Dr, Austin, TX, 78729, USA

\* Corresponding author email: suhaankhanisme@gmail.com

## Abstract

Gravitational lensing has been used for over three decades as effectively the only tool to reveal the nature of "dark" sources in the Milky Way. "Dark" sources include stars, planets, black holes, or even dark matter. We hypothesized that by using 600,000 stars as data obtained from The Zwicky Transient Facility, a telescope at the California Institute of Technology, we could use Artificial Intelligence to increase the speed and accuracy of microlensing detection. The main objective of this project was to distinguish between microlensing and non-microlensing events, the former of which could be dark matter events. Through the rigorous testing of various Artificial Intelligence models, including Linear Regression, KNN, SVC, Decision Tree, and others, we conducted a comparative analysis to discern disparities in accuracy, precision, and other recall among these models. Through hyperparameter tuning and the elimination of false positives, the KNN model was confirmed to be the most accurate model, generating a 97% accuracy for the detection of microlensing events within the Milky Way. This scientific inquiry aims to improve future searches for microlensing events, ultimately expediting the current scientific processes employed in microlensing detection and helping astronomers place search constraints on the nature of dark matter through AI.

## Keywords

Physics and Astronomy, Astronomy and Cosmology, Dark Matter Detection, Gravitational Microlensing, Machine Learning.

## Introduction

For nearly a century, scientists and astronomers alike have wondered what constitutes dark matter. Just a century ago, it was unknown whether there was even another form of matter. Gravitational microlensing has emerged as one of the most prominent methods to detect dark matter, among other dark objects in the Milky Way, yet observing it was once thought impossible. Albert Einstein himself remarked, "There is no hope of observing this [gravitational lens] phenomenon directly".<sup>1</sup> It was previously believed that matter made up everything in the universe until Swiss researcher Fritz Zwicky, who studied galaxy clusters while working at the California Institute of Technology, Pasadena, CA, obtained evidence of unseen mass he called Dunkle Materie (dark matter).<sup>2</sup>

The concept of dark matter was further supported by the work of Vera Rubin and Kent Ford, who observed the rotation curves of galaxies and found that the outer regions rotated much faster than could be accounted for by visible matter alone.<sup>3</sup> This discrepancy suggested the presence of a significant amount of unseen mass, reinforcing the need for methods to detect dark matter.

Scientists today have found that dark matter, imperceptible to our senses, is six times more abundant in the universe than the matter we're acquainted with. This makes it 80% more common, yet scientists know little about it.<sup>4</sup>

Dark matter does not interact with electromagnetic waves, meaning it does not absorb, reflect, or emit any light waves, making it extremely difficult for scientists to detect.<sup>4</sup> Dark matter's elusiveness caused it not to be the subject of intense research until four decades ago when astrophysicist Bohdan Paczynski discovered and shared his findings on gravitational microlensing detection on an isolated object while attempting to find dark matter.<sup>5</sup> He wrote that because dark matter has mass, and all objects with mass bend light, fluctuations in magnitude (light levels) of stars can often be used to detect it. Such events became known as microlensing events.

Ten years later, Charles Alcock et al. expanded this theoretical claim and put it into practice. He reported fluctuations in light waves in objects with a 95% confidence level of dark matter detection.<sup>6</sup> Previous scientists, such as Paczynski, inspired Charles's method. With Charles's method, fluctuations in light curves could be used to detect dark matter. Specifically, a large spike in a light curve would indicate a significant increase in source brightness, highlighting a potential dark matter event. This means that dark matter, while undetectable by other methods, could be detected indirectly by examining its gravitational effects.

The problem, however, was that Charles and other scientists' methods took far too long and were not as accurate as many astronomers desired. Current methods have not changed significantly, and humans are still attempting to analyze photometric light curves with noise and fluctuations, which is time-consuming and sometimes inaccurate. The large amount of noise present in time-series data provides another challenge that causes the graphs of microlensing events and the graphs of false positives, non-microlensing events incorrectly assumed to be microlensing events, to look very similar.

False positives are the most common form of misinterpreted data in current methods, and reducing them would significantly improve accuracy.<sup>7</sup> Since human detection cannot reliably distinguish between microlensing and non-microlensing graphs, non-microlensing graphs are often mistakenly classified as microlensing events. As a result, stars with no microlensing activity have been studied for years for the presence of dark matter, even though no trace may be present.

With the advent of Artificial Intelligence, we hypothesized that integrating AI into the microlensing detection process could increase efficiency and accuracy compared to human-based detection. Recent advancements in AI have shown promise in various astronomical applications, demonstrating the potential for machine learning models to enhance data analysis and improve detection rates.<sup>8</sup> For instance, Machine Learning techniques have been successfully applied to detect exoplanets and classify astronomical objects, indicating a growing trend of utilizing AI in astrophysics.<sup>9</sup> This paper details the data, methodology, results, validations, and analysis of using AI to distinguish between, and hence classify, gravitational microlensing and non-microlensing events. The tested models were the K-Nearest Neighbors (KNN), Decision Tree, Support Vector Classifier (SVC), and Logistic Regression model. Each model presented a high accuracy, yet after hyperparameter tuning, the KNN model was the most accurate, with 97% accuracy.

## Methods

### ***Data from the Zwicky transient facility***

The dataset used during this project was from the Zwicky Transient Facility (ZTF), a telescope operated and maintained by the California Institute of Technology, and located at the Palomar Observatory. Built in 2018, ZTF has an extremely large camera lens that can scan half the northern sky in just one day.<sup>10</sup> It is designed to detect light changes in astronomical objects and provides an extensive dataset for projects that involve detecting fluctuations in light. We sampled > 600,000 images of stars, focusing on key statistics such as Sigma values (10/25), Von Neumann ratios, mean and min/max magnitude, and skewness.

Sigma 10 is a statistical measure in astrophysics, indicating the significance of a detected signal compared to background noise. It is the number of standard deviations from the mean of background noise, with higher sigma values denoting stronger signals and more confidence in detection. Astronomers use sigma thresholds to distinguish real astronomical phenomena from random data fluctuations; a detection at Sigma 10 of > 1 or higher is usually deemed significant in astrophysics. Sigma 25 is similar but has a higher threshold for significance than Sigma 10. Skewness and the Von Neumann Ratio are two other variables that measure the significance of the brightening/dimming of a galactic object. Positive values indicate dimming, and negative values indicate brightening, with larger values showing a more significant effect.

ZTF also includes photometric light curves of each star presented as graphs of time vs. magnitude (the amount of light entering the telescope from the star). This was useful since the AI could be trained on graphs of real microlensing/non-microlensing events and not on simulations that could introduce bias or compromise accuracy. Astronomers have long used fluctuations in light curve graphs to detect changes in magnitude for exoplanet detection and microlensing detection. For a gravitational microlensing event, photometric light curves are examined for sharp increases that indicate the influence of dark matter on the light passing between the star and telescope (in our case, ZTF). While there may be other astronomical phenomena, such as black holes and neutron stars, that can cause a spike, dark matter shows a larger increase in the time vs. magnitude graph because the object is more massive. This is supported by Einstein's Theory of Relativity, as more massive objects tend to bend the most light as they travel through the space-time continuum. More intense spikes are attributes of dark matter, and the AI can be trained on light curves from the ZTF.

### ***Data preprocessing***

The dataset contained other quantitative data or data that needed some specific parameters to be viable, both of which were not specifically necessary for our purpose. To filter out this data, we used preprocessing techniques and Python code in Google Collaboratory to limit the number of relevant observations to < 50, the filter code to red, and only train the model on a dataset of stars, filtering out the galaxies.

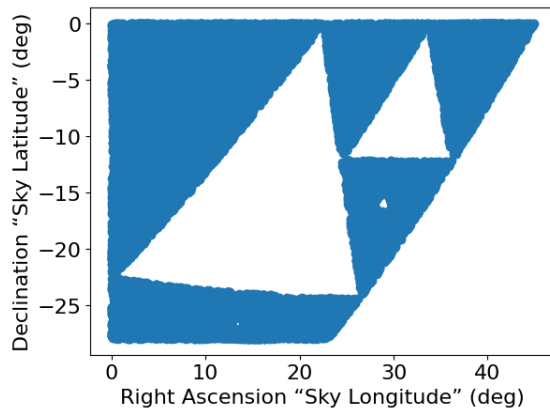
The purpose of filtering the number of relevant observations was to include only stars that have been reported and studied previously to ensure that the data was as accurate as possible and that statistics were as unbiased as possible. The filter code refers to the camera filter on the ZTF when the photographs were taken. There are two camera filters: red and green. These filters capture longer and shorter wavelengths of light, respectively. The higher sensitivity of the red filter helps detect small but significant changes in the light curves. For this reason, we decided to remove data from the green filter and focus on data from the red filter. Finally, differentiating between stars and galaxies is important because only stars

can experience gravitational lensing events that are large enough to reveal dark matter. We further filtered the dataset to galactic objects classified as 'stars' by the ZTF, which left 239,000 star objects to train the AI model on.

## Results and Discussion

### *Data Visualization*

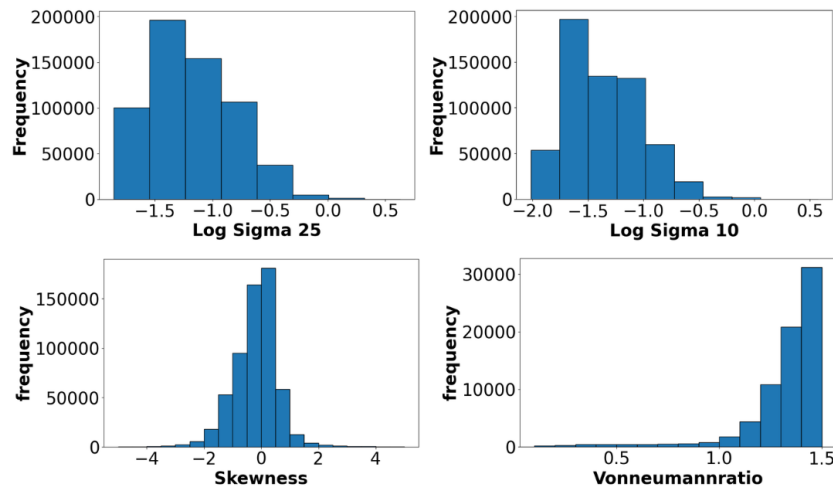
We chose our sample to arbitrarily lie across right ascension and declination coordinates. Declination is the equivalent of latitude and is expressed in degrees, and right ascension is the equivalent of longitude and is also expressed in degrees (Figure 1).



**Figure 1.** Distribution of Galactic Objects Chosen from the ZTF Graph of Ascension and Declination Coordinates for Sample Data. While the sample may seem to have some pattern based on telescope scanning, the coordinates were selected randomly to ensure no bias; this allowed for a wide variety of stars to train AI on. Though the triangles seem intentionally crafted, the selection of stars is random.

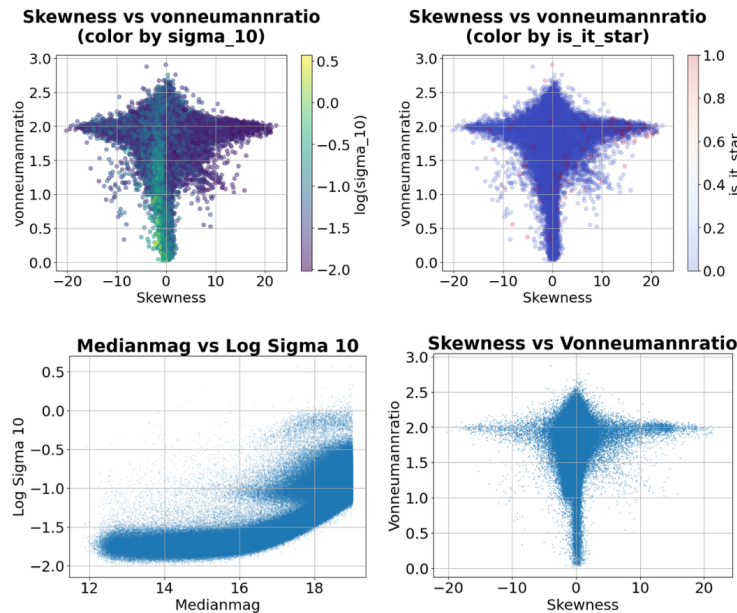
To maximize results, it became crucial to identify which parameters (features) to use for the AI models to be trained on. The larger the association/correlation of the features with the magnitude of the light curves, the more accurate the AI performance would be. Using Matplotlib, we created graphs of the Von Neumann ratio (only values of 0-1.5 were graphed to show the general pattern), (log) sigma values, skewness, and magnitude to understand the ZTF dataset better and identify the patterns to train the AI on (Figure 2). The X-axes of the Sigma 10 and Sigma 25 graphs, plotted as the logarithm of the Sigma values, show the actual Sigma values of the population of > 239,000 star images.

After graphing these images, it became evident that there were significantly more non-microlensing event stars than microlensing event stars. Sigma 10, Sigma 25, and Von Neumann Ratios are not particularly high, supporting this hypothesis. Over 90% of Sigma 10 values are < 0.8 mag (this corresponds to a log value of -0.097), indicating that only a light change with 0.8 mag standard deviations has occurred, which does not indicate a microlensing event. Typical values that do indicate a microlensing event are greater than 1 mag, and high-confidence ones are greater than 2 mag on a linear scale.



**Figure 2.** Distribution of Key Parameters for Potential Microlensing Events. Log Sigma 25 (top left), Skewness (bottom left), Von Neumann Ratio - values of 0-1.5 only shown - (bottom right), and Log Sigma 10 values (top right) are compared against frequency to visualize the data better and understand patterns between different parameters and frequency. Visualizations indicate a wide variety of data from both stars and non-stars. Very few events have a Sigma 10 Value  $> 1$ . Those that are  $> 1$  may be indicative of microlensing events.

To further identify patterns, Sigma 10, Median Magnitude, Skewness, and Von Neumann Ratio were compared against each other to further see trends in the dataset (Figure 3). Comparing them revealed that each had trends that the Artificial Intelligence model could use. For example, the Skewness vs. Von Neumann Ratio graph colored by whether the galactic object is a star indicates that stars, the only objects that can have microlensing events occur, have values of Skewness greater than -5 and values of Von Neumann Ratio greater than 1.5 over 70% of the time.



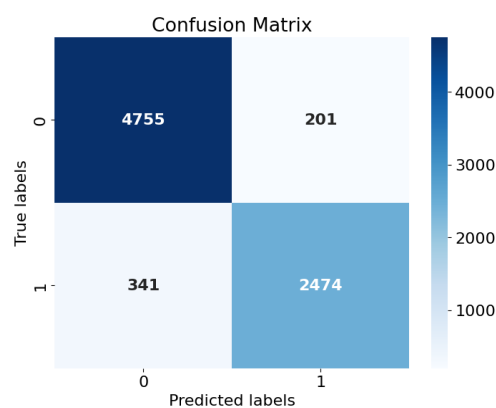
**Figure 3.** Relationships Between Key Microlensing Parameters and Classifications. Median Magnitude vs. (Log) Sigma 10 values (bottom left), Skewness vs. Von Neumann Ratio (bottom right), Skewness vs. Von Neumann Ratio (colored by whether the object identified as a star or not) (top right), and Skewness vs. Von Neumann Ratio values (colored by (Log) Sigma 10 values) (top left) are compared. Graphs suggest a clear relationship between these values, which makes them ideal for models to be trained on because of a definitive pattern.

Subsequently, these values were compared against each other to further discern trends in the dataset (Figure 3). Comparing them, we found that each had trends that the Artificial Intelligence model could use. For example, the Skewness vs. Von Neumann Ratio graph colored by whether the galactic object is a star indicated that stars, the only objects that can have microlensing events occur, have values of Skewness greater than -5 and values of Von Neumann Ratio greater than 1.5 over 70% of the time.

### ***Training and Testing***

The AI model was trained in Python programming and was handled completely by Google Collaboratory. The data was first split into training and testing data. 0.8 (80%) of the data was training data, and 0.2 (20%) of the data was testing data. Training data was purposefully set to be a larger sample of data than testing data so that the AI could learn the most before it was put in a testing environment. After multiple trials with different training and testing data proportions, it was clear that 80/20 was the best proportion for the AI to learn from, as accuracy was highest in terms of the detection of microlensing events.

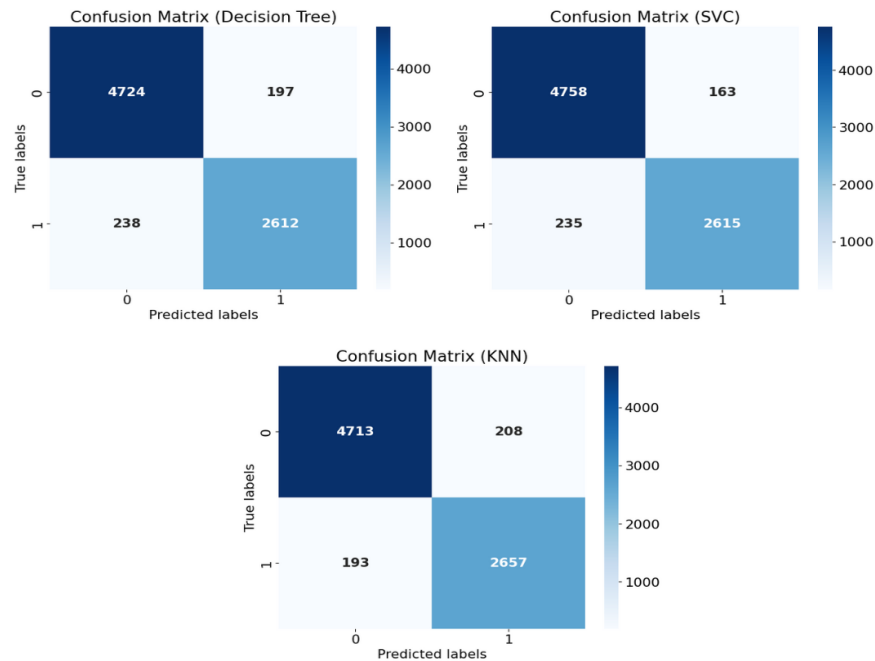
The Logistic Regression model was trained using the sklearn python library, and its confusion matrix was then graphed (Figure 4).



**Figure 4.** Logistic Regression Model Confusion Matrix The confusion matrix presents four key metrics: True Positives (TP = 2474), True Negatives (TN = 4755), False Positives (FP = 201), and False Negatives (FN = 341). The sum of TP and TN (7229) is much larger than the sum of FP and FN (542), suggesting that the model accurately distinguishes between the two classes rather than defaulting to one category. This matrix highlights the results from the Logistic Regression Model, the first tested model, visualized using matplotlib.

The AI had a clear distribution of guesses, which was significant because it exemplified that the AI was not simply guessing one singular option (either 0 or 1). Instead, the model made guesses that were not only accurate but also distributed across possible options, which is another good indicator that the 80/20 proportion was a good option for training/testing data. The AI system was then evaluated using three distinct machine learning models: K-Nearest Neighbors (KNN), Decision Tree, and Support Vector Classifier (SVC). These models were specifically selected for their proven capabilities in detecting anomalies indicative of microlensing events. Each model's performance determined its efficacy in identifying potential microlensing phenomena within the ZTF dataset (Figure 5).

Furthermore, each model's confusion matrix was evaluated with 0s and 1s. A '0' indicated no microlensing event, while a '1' indicated a microlensing event. The left side of each matrix is the "true values" of the model, and the bottom side is the "predicted values." Thus, the top left is the true negative, and the bottom right is the true positive.



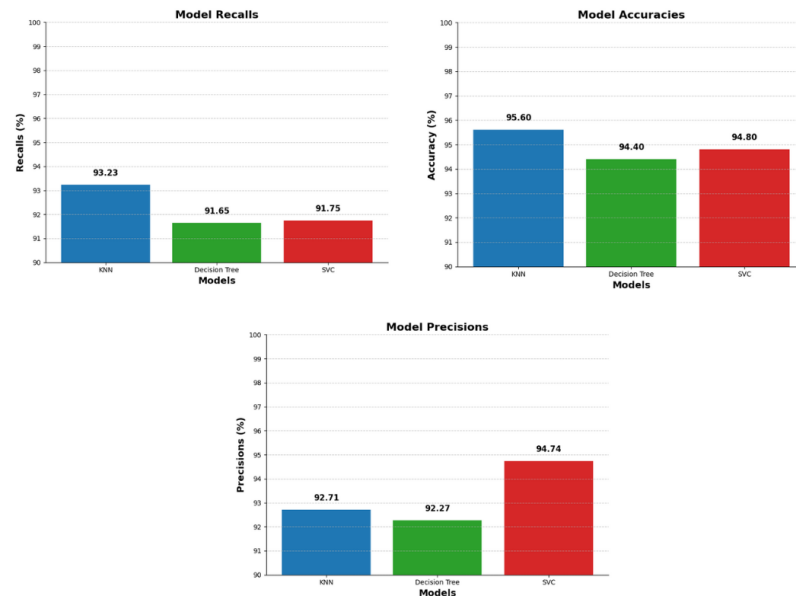
**Figure 5.** Confusion Matrices of Decision Tree, SVC, and KNN Decision Tree (Top left), Support Vector Classifier (Top right), and K-Nearest Neighbors (Bottom) confusion matrices. Matrices indicate that all models have similar guesses in each category and are correct (true positive and false negative) over 90% of the time. The models above are more accurate than Logistic Regression and were, therefore, the primary focus in hyperparameter tuning.

### ***Statistical Analysis Through Accuracy, Precision, Recall***

We used precision, recall, and accuracy, three statistical analysis measures to determine the quality of our results. Accuracy is the amount the model got correct (true positives and true negatives) over the total. For all three models, the accuracy was very similar, with an accuracy of  $96\% \pm 2\%$ .

With such high accuracy for all three models, we resorted to other statistical analysis measures, namely precision and recall. Precision is a metric that measures how well a machine learning model can correctly predict positive instances while minimizing false positives.<sup>10</sup> In our case, this meant maximizing the number of times our model predicted microlensing events when it was a microlensing event, which is the bottom right of the confusion matrix. Precision is essentially (True Positives) / (True Positives + False Positives). The Precision of all three models was also similar, yet there were more noticeable differences than in the case of accuracy. For one, the precision of the SVC model was slightly over 2% (94.74%) higher than the precision of the Decision Tree (92.27%) and KNN models (92.21%). Second, the precision of the model had a higher margin of error than accuracy, with  $95 \pm 0.6\%$  with a 99% confidence level (Figure 6). Lastly, after over 50 runs of the model, it became clear that precision was lower than accuracy over 95% of the time.

While precision focuses on the correctness of positive predictions, recall measures how often a machine learning model correctly identifies positive instances (true positives) from all the actual positive samples in the dataset.<sup>10</sup> The formula to calculate recall is (true positives) / (true positives + false negatives). The Recall for all three models was  $91\% \pm 2$  with a 99% confidence level. Again, the KNN model slightly outperformed the others, averaging 93.15% over 25 trials, while the Decision Tree and SVC models had an average of 91.5% (Figure 6). While this difference may not initially seem substantial, it confirms that the KNN model seems to be the most accurate one thus far.



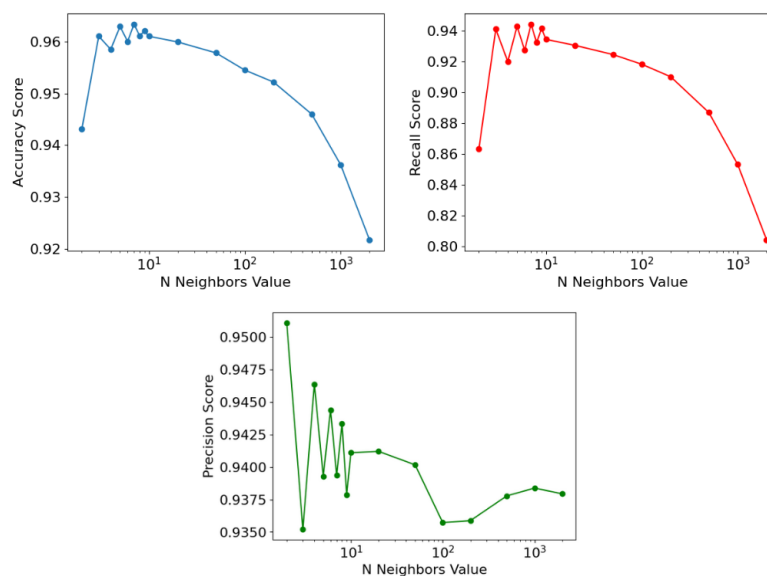
**Figure 6.** Mean Accuracy (Left), Mean Precision (Right), and Mean Recall (Bottom) over 50 trials prior to hyperparameter tuning. On average, all metrics are about equal ( $\pm 2\%$ ) independent of model choice and metric. SVC precision is the only metric that seems to be higher than the rest.

## Hyperparameter Tuning

The high accuracy, precision, and recall observed in our data model indicate that Artificial Intelligence models hold promise for gravitational microlensing detection in the future, facilitating faster analysis compared to human detection methods. The K-Nearest Neighbors (KNN) model emerges as the superior performer among the evaluated statistical measures. Overall, the KNN model exhibited the highest accuracy, recall, and second-highest precision among these models, establishing it as the leading candidate for microlensing detection.

However, improvements can still be made. Now that we have realized that the KNN model seems to perform the best, we used hyperparameter tuning to further refine the model. Hyperparameter Tuning is the process by which an algorithm is created to test for nearly every plausible parameter and see at which value a certain parameter causes the model to perform at the highest level possible.<sup>12</sup> For the KNN model, the parameter often optimized is  $k$ , which is the number of neighbors for the model.<sup>12</sup> Changing  $k$  can lead to underfitting or overfitting as the model uses many nearby samples to make its prediction. When  $k$  is large, the model incorporates a broader range of neighbors, potentially oversimplifying the decision boundary and leading to underfitting by ignoring local patterns and nuances in the data. Conversely, smaller values of  $k$  focus on fewer neighbors, which may capture noise and result in overfitting by fitting too closely to the training data. Finding the optimal  $k$  involves balancing these factors to achieve a model that generalizes well to new data while accurately capturing the underlying relationships within the dataset.

To achieve this goal, we built an algorithm using Python code that tests for various  $k$ -values, from 1-1000, reporting the accuracy, precision, and recall. We then generated graphs for each analysis method to find the best  $k$ -value for the KNN model. With hyperparameter tuning, the best  $k$ -value for the KNN model was seven neighbors, resulting in an average accuracy of 96% (Figure 7).



**Figure 7.** Accuracy, Precision, and Recall graphs with different experimental values of  $N\_Neighbors$ . Hyperparameter Tuning is used to run a for loop in which 10 values between 1

and 1000 were tested. N\_Neighbors value of 7 has the highest average of precision, recall, and accuracy and was used to optimize the model.

Unfortunately, no single value optimized all three statistical measures. The highest precision was achieved with  $k = 1$ , and the highest recall with  $k = 6$ . However, the K-value of 7 provided the best balance of precision, recall, and accuracy. With this high accuracy, this AI model's use to detect stars with potential microlensing events becomes even more useful to the scientific community.

To validate the results, we compared the AI model to human researchers in the field. The general feedback suggests that data efficiency is the biggest benefit of integrating artificial intelligence into this predominantly human domain. This model can significantly increase the speed at which we analyze microlensing graphs. In fact, researchers currently in the field report that this model can accelerate the dark matter detection process by over tenfold while maintaining or even slightly increasing accuracy compared to human researchers.

## Conclusion

While the initial results seem promising, the model still needs further testing in practical applications. A follow-up study will compare human researchers directly against the AI model to determine whether there are differences in speed and accuracy and to what extent those differences actually lie.

Despite these potential advantages, it is important to acknowledge the model's limitations. One major limitation was that the model was only trained on the ZTF database. While it contains some of the most expansive and useful data for this project, more distinct data from different databases would create more comprehensive, balanced, and accurate results for the model. To counter these limitations, we added data from the European Space Agency (ESA) mission Gaia database, which has a census of stars, including their brightness, temperature, composition, and evolutionary stage. This database was combined with ZTF to provide even more variety. Combining two datasets provided a larger variety and refined the models by reducing the bias that would naturally arise from using just one dataset. Future studies should similarly use combined or mixed datasets to reduce bias and achieve greater variability, further refining Artificial Intelligence models. When the results were run again with the combined database, as hypothesized, all models performed at a slightly higher accuracy (+2.4% on average), though precision and recall did not change. The KNN model was still the most accurate, extending its lead to 3% with an average accuracy of 97%. This accuracy could be further improved by creating thousands of simulations of microlensing events graphs with noise to train the AI to have a larger data set. However, this could also decrease the true accuracy, as simulations are never as precise as true graphs of microlensing events, and factors such as noise can be extremely hard to duplicate.

The next step is to detect dark matter directly. Currently, the AI model can differentiate between microlensing and non-microlensing events. However, that does not guarantee that a dark matter event has occurred.

In fact, there can be other reasons for microlensing events than dark matter, such as Black Holes and Neutron Stars, which do not suggest traces of dark matter but still bend the light, indicating a microlensing event. Therefore, the next step is to filter and distinguish astronomical objects that do not indicate dark matter from those that do. Taken together with this study of microlensing detection, this work may reveal new insights into the formation, composition, and evolution of dark matter and the Big Bang Theory.

## Acknowledgments

Thanks to Tony Rodriguez for providing feedback to further this research project. This research would be incomplete without their help. This work has made use of the Zwicky Transient Facility operated by Caltech as well as the Gaia Database<sup>13</sup> operated by the ESA. Matplotlib<sup>14</sup>, Numpy<sup>15</sup>, Sklearn<sup>16</sup>, and Pandas<sup>17</sup> were all imported libraries used. All code was written in Python in Google Collaboratory.

## References

1. Einstein, A. Lens-Like Action of a Star by the Deviation of Light in the Gravitational Field. *Science* 1936, 84, 506–507. <https://doi.org/10.1126/science.84.2188.506>.
2. Bernstein, M. Theorists Imagine a Different Kind of Dark Matter | Symmetry Magazine. *Symmetry Magazine*, Jan 4, 2022. [www.symmetrismagazine.org/article/theorists-imagine-a-different-kind-of-dark-matter?language\\_content\\_entity=und](https://www.symmetrismagazine.org/article/theorists-imagine-a-different-kind-of-dark-matter?language_content_entity=und).
3. Rubin, V. C.; Ford, K. W. Jr. Rotation of the Andromeda Nebula from a Spectroscopic Survey of Emission Regions. *Astrophys. J.* 1970, 159, 379. <https://doi.org/10.1086/150317>.
4. Smulsky, J. Dark Matter and Gravitational Waves. *Nat. Sci.* 2021, 13 (3), 76–87. <https://doi.org/10.4236/ns.2021.133007>.
5. Paczynski, B. Gravitational Microlensing by the Galactic Halo. *Astrophys. J.* 1986, 304, 1. <https://doi.org/10.1086/164140>.
6. Alcock, C.; Allsman, R. A.; Alves, D.; Axelrod, T. S.; Becker, A.; Bennett, D. P.; et al. The MACHO Project: Limits on Planetary Mass Dark Matter in the Galactic Halo from Gravitational Microlensing. *Astrophys. J.* 1996, 471, 774–782. <https://doi.org/10.1086/177999>.
7. Cowan, P.; Bond, I. A.; Reyes, N. H. Towards Asteroid Detection in Microlensing Surveys with Deep Learning. *Astron. Comput.* 2023, 42 (1), 100693. <https://doi.org/10.1016/j.ascom.2023.100693>.
8. Baron, D. Machine Learning in Astronomy: A Practical Overview. *arXiv Preprint* 2019, arXiv:1904.07248. <https://doi.org/10.48550/arXiv.1904.07248>.
9. Wyrzykowski, Ł.; Rynkiewicz, A. E.; Skowron, J.; Kozłowski, S.; Udalski, A.; Szymanski, M. K.; et al. OGLE-III Microlensing Events and the Structure of the Galactic Bulge. *Astrophys. J. Suppl. Ser.* 2015, 216 (1), 12. <https://doi.org/10.1088/0067-0049/216/1/12>.
10. Kostadinova, I. ZTF. *Zwicky Transient Facility*. [www.ztf.caltech.edu](http://www.ztf.caltech.edu) [Accessed Dec 21, 2024].
11. Evidently AI. Accuracy vs. Precision vs. Recall in Machine Learning: What's the Difference? *Evidently AI*. [www.evidentlyai.com/classification-metrics/accuracy-precision-recall](https://www.evidentlyai.com/classification-metrics/accuracy-precision-recall) [Accessed Dec 21, 2024].
12. Feurer, M.; Hutter, F. Hyperparameter Optimization. In *Automated Machine Learning: Methods, Systems, Challenges*; Springer: Berlin, 2019; pp 3–33. [https://doi.org/10.1007/978-3-030-05318-5\\_1](https://doi.org/10.1007/978-3-030-05318-5_1).
13. Agrawal, S. Hyperparameter Tuning of KNN Classifier. *Medium*, June 4, 2023. <https://medium.com/@agrawalsam1997/hyperparameter-tuning-of-knn-classifier-a32f31af25c7>.
14. Hunter, J. D. Matplotlib: A 2D Graphics Environment. *Comput. Sci. Eng.* 2007, 9, 90–95. <https://doi.org/10.1109/MCSE.2007.55>.
15. Harris, C. R.; Millman, K. J.; van der Walt, S. J.; et al. Array Programming with NumPy. *Nature* 2020, 585, 357–362. <https://doi.org/10.1038/s41586-020-2649-2>.
16. Pedregosa, F.; Varoquaux, G.; Gramfort, A.; Michel, V.; Thirion, B.; Grisel, O.; et al. Scikit-Learn: Machine Learning in Python. *J. Mach. Learn. Res.* 2011, 12 (12), 2825–2830. [www.jmlr.org/papers/volume12/pedregosa11a/pedregosa11a.pdf](http://www.jmlr.org/papers/volume12/pedregosa11a/pedregosa11a.pdf).

17. McKinney, W. Data Structures for Statistical Computing in Python. In *Proceedings of the 9th Python in Science Conference*; van der Walt, S., Millman, J., Eds.; 2010; pp 51–56. <https://doi.org/10.25080/Majora-92bf1922-00a>.

## Authors

Suhaan Khan is a senior studying in Austin, Texas. He is passionate about Astrophysics and AI research and plans to major in Artificial Intelligence. Suhaan has various experiences at State, National, and international research conferences and science fair competitions. He hopes he can use his knowledge of AI to create a lasting impact on society.

## Reviewer Suggestions

IJHSR requires authors to recommend 3 - 6 qualified reviewers in the same discipline the research is conducted in.

Please provide the first and last name, institution, city/state/country of institution, and email address for each suggested reviewer.

Recommended Reviewers must:

- Be fluent in English
- Be from different institutions (limited to 1 person per institution)
- Have not collaborated with the author's work
- Not be affiliated with the author's institution
- Should have PhD in the research area or more than ten years of experience in the field

| First Name and Last Name | Institution | City/State/Country | Email |
|--------------------------|-------------|--------------------|-------|
|--------------------------|-------------|--------------------|-------|

1.

Laura Hermosa Muñoz | Institute of Astrophysics of Andalusia | Granada, Spain  
lhermosa@cab.inta-csic.es

2.

Florian Peißker | First Physics Institute University of Cologne | Köln, Germany  
peissker@ph1.uni-koeln.de

3.

Hai-Long Huang | School of Fundamental Physics and Mathematical Sciences, Hangzhou Institute for Advanced Study | Hangzhou, China | huanghailong18@mailsucas.ac.cn

4.

Dr. Jillian M. Scudder | Oberlin College | Oberlin, Ohio | jillian.scudder@oberlin.edu