

Machine Learning in Adolescent and Young Adult Mental Health Prediction: an Investigation

Haowen Yang

*Johnbapst, Maine, USA
hyang28@johnbapst.org*

Abstract. Adolescent and young adult mental health is a pressing global public health concern—recent data shows ~20% develop at least one mental health disorder by age 24. Traditional assessments (e.g., PHQ-9, GAD-7) rely on self-reports and professional judgment, leading to subjectivity and underreporting due to stigma, which hinders large-scale early screening. In the view, machine learning offers a practical solution amid the advancements of AI and big data. This review covers key approaches: traditional models (decision trees, random forests, SVMs) and deep learning models (ANNs, CNNs, RNNs), which process diverse data (social media content, wearable sensor data, speech signals) to uncover complex patterns beyond human observation. Markedly, there are still critical issues such as model opacities (black boxes), ineffective cross-cultural extrapolation, and privacy. These problems are dealt with by emerging technologies such as explainable AI (XAI), transfer learning, and federated learning. Machine learning shows a lot of promise in early mental health detection, but requires better transparency, reliability, and ethical usage to be widely used in clinical practice.

Keywords: Machine learning, mental health prediction, deep learning

1. Introduction

During the last decade, it can be observed that there is a definite tendency to increase the number of teenagers and young adults with mental issues- in particular, college students. Many of them are struggling with chronic stress, anxiety, and depression, fueled by heavy academic demands, intense competition to get into college and secure employment, family pressures, relationship issues, and overuse of social media. These pressures do not only affect emotional stability, but also cognitive abilities such as attention and memory. Even better, they frequently result in a decline in academic achievement, increased dropouts, and even risky behaviors like drug abuse.

In most extreme situations, youths fail to realize that they are experiencing some mental health problems until the symptoms of those problems affect their lives- such as skipping classes, avoiding friends, or feeling overwhelmed with ordinary chores. This is why it is so important to identify it at an early age. Until some stress-inducing situations come by, clinical mental health issues need immediate treatment, and early detection of warning signs can significantly optimize the results of young individuals.

Mental health assessment is highly dependent on structured questionnaires and clinical interviews in schools and healthcare systems. Depression and anxiety assessment instruments such as PHQ-9 and GAD-7 are user-friendly and cost-effective, not to mention easy to administer [1]. However, in the experience of working with clinical teams, such tools are obviously flawed. To begin with, they rely wholly on youths telling the truth about their symptoms. Some are afraid of being judged or called mentally ill, other just do not even realize that their stress can be the sign of the bigger problem.

Besides, the personal bias of evaluators may cause the lack of consistency. It can be witnessed situations where one counselor would categorize the reaction of a student as having mild anxiety, and another categorizes it as being within the normal range. This approach also takes a lot of time and manpower- so it cannot be applied in large groups such as a whole high school or university campus. It is obviously required more effective and objective methods to filter mental health dangers, particularly among the youth who could fall between the cracks with the conventional methods.

Recent researchers make use of various sources of data to train machine learning prediction models [2, 3]: social media posts (De Choudhury et al. [4]), lifestyle metrics (Dham et al. [5]), clinical interview recording (Alhanai et al. [6]) and wearable sensor data (Tsai et al. [7], Mullick et al. [8]). Although these studies are promising in the real sense, other critical issues such as model transparency and cross cultural adaptability remain unresolved before such tools can be extensively relied upon in clinical practice.

2. Machine learning in mental health prediction

2.1. Conventional machine learning approaches

Conventional machine learning (decision trees, random forests, SVMs) are a balance between performance and usability. They need fewer computational resources and are more transparent, which is important in gaining the trust of clinicians. These approaches are the standard of low-resource clinics and initial-stage research, as this paper reviewed the relevant literature.

2.1.1. Random forest

Random forest is the most popular ensemble approach, and it is better at mixed-type data (such as sleep logs, stress scores, and lifestyle surveys) which combines the predictions of multiple decision trees. This design removes the personal bias and the possibility of overfitting, which is often the case in which models do well on the training data but do poorly in practice.

In a study by Tsai [7], random forest was trained to predict panic attacks in young adults with a history of anxiety, with wearable data as the input to the model: heart rate variability (one of the most important indicators of nervous system operation), efficiency of sleep, and the number of steps taken daily. The team identified the most effective features before the training- heart rate variations in sleep, sleep efficiency of less than 70 percent, and afternoon activity peak- during the training. The findings were striking: the accuracy was between 67.4% and 81.3, the sensitivity was high (78.2%), as well as the specificity (75.6%). However, the most important, the model predicted warning signs 1-3 days prior to a panic attack, providing clinicians the opportunity to intervene with coping mechanisms or check-ins.

In a study by Dham [5], random forest was used on the data of 1, 800 college students, which included BMI, GPA, time spent exercising (less than 150 minutes/week was considered to be low)

per day, number of hours sitting (>8 hrs/day) per day, and perceived stress scores. The trends identified in the model were clear: low exercising students that spend more than 20 hours sitting and have stress scores higher than 20 were three times more likely to report depressive symptoms. The interaction effect was surprising to me, as regular exercise reduced the adverse effect of high academic stress, which is a protective effect of exercise on mental health.

Other weaknesses of random forest are poor interpretability of feature combinations. Although it is possible to discover the features that are important (such as sleep efficiency), it is not possible to easily follow how the model interacts between those features to arrive at a prediction. Some clinicians are not willing to use this black box like problem when making high-stakes decisions.

2.1.2. Decision tree

There is a high likelihood that decision trees are the most intuitive machine learning tool to be used in mental health. This paper tends to describe them to clinicians as digital flowcharts- they organize data through simple rules that are step by step and whose rules are based on explicit thresholds. As an example: first, does the student have less than 6 hours of sleep a night? When yes, proceed on to the next branch: is their perceived stress score over 22? In that case, put them at high risk of depression.

In a study by Dham [5], the analysis of the data about university students was conducted with the help of decision trees, and the conclusions were rather simple: students who were physically active (at least 3 days/week of moderate-vigorous exercise) always scored lower in the scales on depression and anxiety. Conversely, individuals with GPAs of less than 2.5 (based on a 4.0 scale), less than 2 days of physical activity per week, and more than 9 hours of non-physical activity per day were much more likely to demonstrate clinically significant symptoms of depressive disease. The tree also identified distinct risk levels, including sleeping less than 5.5 hours a night, doubled the risk of anxiety, and a perceived stress score over 22, which tripled the depression risk.

The best part of it is that decision trees enable clinicians and the youth. Once a student is identified as high risk, the tree will show precisely what factors contributed to that outcome- such as poor sleep and lack of exercise. This is not merely a yes/no tagline, it is a change map. It can be observed clinicians utilize these trees to engage in a conversation with someone such as, the risk score is high because sleeping 5 hours a night and not exercising. Let beginning with some small changes-goal is 6.5 hours of sleep and 20 minutes of walking a day- and see how feel.

But decision trees are not without their weaknesses. In the experience, in practice, they are too sensitive to minor changes in data, and even a few entries changed can change the entire tree structure and give inconsistencies. They also tend to overfit particularly when using deep, complex trees. This is why they are commonly pruned down or pooled together in ensemble models such as random forests to enhance reliability in practice.

2.1.3. Support vector machine (SVM)

SVM are especially effective with high-dimensional data, such as text, images, or genetic data, and have become a popular option to analyze unstructured data such as posts on social media in mental health studies. The difference is that they can flourish with more features (such as individual words or phrases) than samples (such as users), a typical situation with social media data.

It can be observed that scientists tend to use SVM more in the study of linguistics due to the abundance of subtle information in the social media posts: word usage, structure of the sentence, and tone of feelings. Such cues are not noticeable to humans, and SVM is very good at identifying trends

that are related to mental health problems. The example of the study conducted by De Choudhury [4] is quite impressive: the authors analyzed more than 100, 000 social media posts of 2, 000 young adults (18-24 years old), and retrieved such characteristics as emotional tone, sentence length, use of pronouns and focus on the topic. The model was trained using 70 percent of the data and tested using the remaining 30 per cent of the data with an accuracy of 72 per cent, sensitivity of 68 percent, and specificity of 75 percent.

The most notable result in the opinion was the way linguistic patterns changed based on the status of mental health: the use of negative words such as sad, lonely, or hopeless were more frequent by young adults with depression, negative words were more often used in the first person (indicating self-focus), and social interactions were rarely mentioned. There are severe limitations of SVM. Its effectiveness plummets drastically in other cultures or languages-linguistic communication of depression in English is not easily transferred to Spanish or Mandarin, etc. Also, the process of tuning the SVM kernel functions and hyperparameters is time-consuming, which may be a limitation to small research teams or small clinics. SVM is also not very efficient with noise in real world data such as sarcasm or language variations which may disturb predictions.

2.2. Deep learning methods

Deep learning models (ANNs, CNNs, RNNs) have revolutionized the mental health prediction task by addressing the types of data that can not be processed by the traditional methods, such as text, audio, images, and complex physiological signals. The best part is that these models emulate human cognition: they do not necessarily follow pre-determined rules, but learn to discern subtle patterns themselves. This is ideal in mental health, where the cause of conditions is often not straight forward- an integration of biological, psychological and social factors.

2.2.1. Artificial neural network (ANN)

Deep learning is based on artificial neural networks (ANN) that are designed to resemble the interlinked neurons of the human brain. This paper thinks this is their best strength as they can capture complex interactions between variables. Mental health data is seldom clear-cut, such as the risk of depression in a student due to poor sleep, high stress, and low social support as a combination. ANN is particularly good at undoing these relationships, which decision trees or SVMs are not.

This was perfectly illustrated in the study conducted by Dhams team [5] where they compared ANN with traditional models. They input the ANN with 12 features of BMI, GPA, sleep duration, exercise frequency, academic stress, social support, and others of more than 1, 800 college students. The outcome was self-explanatory: ANN scored 78% accuracy, being better than decision trees, SVM, and random forests. The researchers explained this success with the fact that the model was able to identify the subtle interactions such as the fact that regular exercise cushions the fact that high academic pressure affects the individual or that sufficient sleep cushions the effect of social isolation.

But ANN has serious disadvantages that restrict application in the real world. To begin with, it requires huge amounts of data to be effective. It can be observed small-scale experiments where ANN does not perform as well as traditional models due to under-fitting: the training data is insufficient to train the ANN, thus it overfits the training data rather than generalizable patterns. Second, the black box problem is critical in this case. It is not possible to follow the path that the model took to make a prediction like with decision trees. A clinician could score high risk of a

patient, but with no idea of what factors contributed to it- difficult to prescribe specific interventions.

It can be also observed that ANN is very sensitive to data quality. The predictions of the model will be erroneous in case the input data is noisy, incomplete, or biased. To illustrate, when a dataset is not representative of low-income students, the ANN will not be as accurate on this group. This brings up ethical issues: biased models may only worsen the health disparities, and marginalized young people will be under-served. With that said, the potential of ANN can hardly be denied. It can be witnessed how it has been used to forecast depression using fMRI data, discover genetic associations with anxiety, and even custom-tailor treatment regimens to individual responses, which would be unattainable with conventional methods.

2.2.2. Convolutional neural network (CNN)

Convolutional neural networks (CNN) began as tools to process images, however, they have proved to be a game-changer in the mental health research, particularly in speech and audio analysis. The most interesting aspect is how they convert sound into actionable insights. Voice captures subtleties of emotion: changes in pitch, speed of speech, pauses, and tone changes. They are usually subconscious- an individual who is depressed may talk slower without noticing- and CNN is good at identifying them.

The practical implementation is as follows: scientists are turning voice recordings into spectrograms, which are images of sound over time, and CNN layers are used to identify patterns. This is an automated system that eliminates human bias. This paper experienced clinicians missing small changes in speech, yet CNN is always able to find the indicators such as 1.2 second pause (vs. 0.6 seconds in non-depressed peers) or flat tone (little change in pitch) that indicate emotional distress.

The perfect example can be found in the study conducted by Alhanai [6]: the authors examined the audio of clinical interviews of 500 young adults (18 -25), turning the records into spectrograms to study pitch, speech rate, and pause duration. The findings were also consistent with the real-life data; the depressed participants would talk at a rate of approximately 40 words per minute, longer pauses, and the tone variation was minimal. The CNN model achieved 83% accuracy, which is 11% higher than the manual analysis. But it is not watertight. Low-quality audio or noises distract outcomes. It can be observed experiments in which samples recorded at the cafe caused the model to lose accuracy by 20%.

Computational demand is another obstacle. The implementation of CNNs is labor intensive and needs a high-performance hardware, which is not always available to small clinics or under-resourced research groups. This study talked to researchers in third world countries who are unable to recreate these studies due to lack of access to expensive GPUs. Then the problem of why: CNN can explain that a certain pattern of speech is a sign of depression, not why a particular pattern of speech is associated with depression. This background would allow clinicians to develop trust with patients, so something like, the speech is slower, isn't as effective as, Slower speech is a common symptom of depression, and low energy.

Nonetheless the range of CNN is impressive. It can be observed it being used with social media text (the detection of suicide risk based on the post content) and facial expressions (the detection of anxiety based on tight jawlines or furrowed brows). In a study, CNN examined video clips of clinical interviews and identified 89 percent of high-risk teens beating human assessers by 14 percent. The trick here is that CNN is to be used in conjunction with explainable tools: when CNN

identifies a speech pattern, promptly respond with "This is similar to 78% of depressed patients in the data-set.

2.2.3. Recurrent neural network (RNN)

Recurrent neural networks (RNN) are constructed on sequential data, or information that evolves with time, and that is precisely how mental health conditions emerge. RNN is used to follow changes, as opposed to cross-sectional data (a single snapshot): 4 weeks of screen time increase, 2 months of decreased sleep, or a gradual reduction in social activities. These minor, gradual changes are difficult to detect, yet RNN can relate them to the new risks.

RNN was usually defined to clinicians as mental health storytellers. "The process by which an adolescent becomes depressed is not usually a one-day process. They may begin to sleep an hour earlier, and then spend more time on their phone, rather than spending time with their friends, and then they will feel constantly sad. RNN also stores event of the past time steps- correlating these sleep variations with future mood variations. This is why it is best when analyzing longitudinal data of wearables, smartphones, or electronic health records (EHRs) data, which describes a story, rather than a single point in time.

The power of RNN is truly hit home by the work by Mullick [8], where the authors followed 300 adolescents (1318) over 6 months with smartphones and wearables and gathered information on the length of sleep, number of daily steps, screen time, and social networking. To relate changes in daily behavior with depressive symptoms, they trained an RNN with LSTM units, which is essential to prevent the vanishing gradient problem. The results were impressive: 2-4 weeks of screen time more than 7 hours a day, and less than 7 hours of sleep at night, were always followed by the observed mood deterioration.

The RNN achieved an accuracy of 81 percent, and it was more accurate than single-day models of data (65 percent) and random forests (73 percent). In reality, however, RNN has enormous weaknesses. It requires regular, full information, skip a week of sleep tracking or spotty screen time logs, and the predictions get less powerful. It can be observed research in which 10% missing data reduce accuracy by 15-percent. Training is also time consuming particularly with long sequences of time. And, similar to other deep learning models, it is opaque: it is hard to know whether screen time or sleep loss was the more influential factor of a risk score. This complicates the ability of the clinicians to target interventions- clinicians are aware that a teen is at risk, they just do not know what habit to address initially.

In the experience working with a rural clinic, the data dependency of RNN is even higher in under-served regions. A large number of rural adolescents lack the regular use of wearables and smartphones, which results in missing data. This implies that the predictions of RNN are less accurate in such groups- increasing the already existing health inequalities. Future RNN models should be more adaptable, as missing data should not be lost in accuracy.,

3. Discussion

3.1. Challenges

Having reviewed dozens of studies and collaborated with clinical teams, It can be observed three fundamental issues that prevent the implementation of machine learning in the vast majority of clinical settings in young adult mental health. The largest one is transparency- what this paper refers to as the black box crisis. CNN, RNN, ANN deep learning models generate the correct risk scores,

and no one can precisely determine how they reached them. A clinician may report a teenager as at risk of depression but they have no idea whether it is due to sleeping pattern, social media posts or anything different altogether.

This veil destroys confidence. Young people and their families do not merely desire to be labeled as being at high risk, they want to know why. Explainable AI systems such as SHAP or LIME are useful as they can point to the most important features, but they are too difficult to use by busy clinicians. Most models are now focused on accuracy rather than explainability, which has them languishing in research laboratories rather than clinical settings.

The second critical problem is the lack of generalization between groups. The majority of models are trained on Western, educated and urban young adults- think U. S. college students or European teens. However, mental health expression is as diverse as cultures, ethnicity and socioeconomic status. A model that identifies negative posts on social media among U. S. teens may fail to detect depression indicators in rural Chinese teens, who tend to communicate negativity indirectly (they say they are tired, rather than sad), because of cultural expectations. Urban vs. rural differences are also important even within countries, with rural teens less exposed to technology, more distinct lifestyle behavior, and less stigma surrounding mental health, which all confuse models that were trained on city data. This leads to health inequities: underserved populations receive inaccurate predictions, and their needs are not met.

The third roadblock is privacy and bias, which is a roadblock that gets easily dismissed in technical discussions. Mental health data is incredibly personal: recordings of conversations, data of wearable sensors, social media posts all of these demonstrate the intimate aspects of the life of a young person. Their mental health status may be revealed as a result of a data breach, resulting in stigma, discrimination or even damage to their academic or professional opportunities. Federated learning is a solution that provides the possibility to train models on a local device rather than centralize data, but it is not common. This paper interviewed tech teams that estimate federated learning is too costly and complicated in technical terms to implement in small clinics.

Prejudice, in the meantime, infiltrates at all levels. When training data is not representative of Black, Latinx, or low-income teen, this model will not be as accurate in those groups. This paper has watched research in which models falsely label 30 percent of low-income teens as low risk, when they are in fact struggling- this underdiagnosis further increases mental health inequities. Even researchers with the best intentions may overlook this: a sample of young adults may consist primarily of college students, excluding those who work full-time or cannot afford college. Top that off with underserved clinics having no tech resources to support these models, not having the hardware or human skill to operate them, and having a system that serves only to expand the divide rather than bridge it.

3.2. Future directions

The future work in order to transform machine learning in the laboratory into clinics requires prioritizing three topics, with the first one being multimodal data fusion. There is no single source of data that can explain everything. A teenager may not be leaving negative posts online but the wearable will record abnormal sleep patterns and the speech will be slower-paced. Integrating social media text, wearable data, speech signals, and survey data, models can develop a comprehensive perspective related to mental health. This is supported by the review of Kadirvelu [9]: the multimodal models have an average of 82% accuracy whereas single-source models have an average of 73% accuracy.

The most exciting thing to me is that attention-based fusion methods are changing, in that they focus on the cues that are the most relevant. To illustrate it, in the case when the decrease in physical activity is accompanied by the negative tone of voice, the model will showcase these two variables rather than be disrupted by insignificant information. This does not only enhance accuracy but also enhances cross-cultural predictions. Models based on several types of data decrease reliance on culture-data (such as direct negative language) and give attention to universal signals (such as sleep or activity alterations). It can be observed preliminary investigations using cross-cultural data in which multimodal models outcompeted single-source ones by 18 percent of cross-cultural data - a massive leap in health equity.

The second priority is to make AI explainable, really explainable, not simple the addition of complex tools [10]. Clinicians and youths do not require waterfall graphs or technical terms, they require straightforward, practical know-how. The transparency that creates trust could be a model saying, the risk score is high because of 4 weeks of poor sleep and more screen time - let's start with sleep hygiene tips. This paper has collaborated with clinical teams on explainable output co-design, and the change can be seen in the fact that a patient who is aware of the reason is twice as likely to participate in interventions. Next-generation XAI systems should be designed as a collaboration with clinicians, rather than only data scientists- with a focus on speed, readability, and applicability over technical perfection.

Third, it is important to have bigger and more varied datasets, those that capture the entire gamut of the lives of the youth. There are too many existing datasets that are biased towards advantaged groups: college-educated, urban, and wealthy. To address this, researchers are advised to collaborate with community groups, rural health facilities and low-income educational institutions to gather data on underrepresented populations. Here can also be useful federated learning, which does not centralize data but collects it locally, preserving privacy but enhancing diversity. It can be observed pilot projects in which this strategy raised the percentage of low-income teens in a dataset, 12 to 38 percent, and model accuracy by 24 percent among the low-income teens.

Another important step is the integration of machine learning into daily appliances, i.e. smartphones, smart wearables, even smart watches. This will enable models to track mental health in real time, not only in yearly screenings. As an illustration, a bracelet might identify abnormal sleeping patterns and encourage the user to engage in mindfulness or a phone application might tell them they are using social media more than usual and encourage them to take a walk with their friends. But there can be no negotiation with ethical guardrails. AI is not to substitute human clinicians, it is to assist them. This paper has talked to teens who fear that apps can monitor them, and so the openness of data usage and an option to decline it is paramount to establishing trust.

Lastly, it is imperative to deal with practical obstacles. The gap between research and practice will be bridged by developing inexpensive and easy to use tools to be used in low resource clinics, as well as by training clinicians on the use of such models. The participation of stakeholders, including young people, clinicians, ethicists, policy makers in the development of the model will also make sure that these tools can be related to the needs and values in real life.

4. Conclusion

This paper would like to share the conclusion of the research, having reviewed dozens of studies and worked with clinical and research teams: machine learning is a potent mental health prediction tool in adolescents and young adults, but its success will rely on its ability to balance the technical capacity with the human needs. Conventional algorithms such as decision trees, random forests, and SVMs are still important they are easy to understand, transparent, and can be used in low-resource

settings. Deep learning models, in turn, push the boundaries: ANNs unravel the interaction between variables, CNNs decode speech and audio signals, and RNNs monitor changes in behavior over a long period of time. Collectively, these tools provide an objective, scalable alternative to conventional assessments which have never existed before.

The most striking fact, however, is that machine learning is strongest in what it does best, which is processing large, diverse datasets, but this is the greatest challenge as well. The black box issue, lack of cross-cultural generalization, risk of privacy and bias are all the result of the data collection process and model design. The answer is not simply a superior technology but more inclusive research. It is important to see the involvement of youth, clinicians, ethicists, and community leaders in all the processes, including design of dataset, and deployment of models.

The best application of machine learning is undoubtedly in the field of early intervention. Finding early signs of problems, such as 4 weeks of insomnia or slowness in speech, before mental health problems worsen can transform lives. Together with clinical skills and humanistic care, such models can decrease stigma, increase access to help, and provide more young people with the help they need. Mental health prediction does not involve the elimination of human connection in the future; it involves the use of technology to make it better. As the emphasis on transparency, fairness, and user-centric design further develops, machine learning will become an irreplaceable means of promoting the mental health of adolescents and young adults across the globe.

References

- [1] Kroenke, K., Spitzer, R. L., & Williams, J. B. W. (2001). The PHQ-9: Validity of a brief depression severity measure. *Journal of General Internal Medicine*, 16(9), 606–613.
- [2] Mitchell, T. M. (1997). *Machine learning*. New York, NY: McGraw-Hill.
- [3] Jordan, M. I., & Mitchell, T. M. (2015). Machine learning: Trends, perspectives, and prospects. *Science*, 349(6245), 255–260.
- [4] De Choudhury, M., Counts, S., & Horvitz, E. (2013). Predicting depression via social media. *Proceedings of the International AAAI Conference on Web and Social Media*, 7(1), 128–137.
- [5] Dham, P., Peden-McAlpine, C., Morelli, C., et al. (2023). Prediction of mental well-being using machine learning and lifestyle factors in university students. *Healthcare*, 11(11), 1601.
- [6] Alhanai, T., Ghassemi, M., & Glass, J. (2018). Detecting depression with audio/text sequence modeling of interviews. In *Proceedings of Interspeech* (pp. 1716–1720).
- [7] Tsai, C. H., Chen, P. C., Liu, D. S., et al. (2022). Panic attack prediction using wearable devices and machine learning: Development and cohort study. *JMIR Medical Informatics*, 10(2), e33063.
- [8] Mullick, T., Radovic, A., Shaaban, S., & Doryab, A. (2022). Predicting depression in adolescents using mobile and wearable sensors: Multimodal machine learning-based exploratory study. *JMIR Formative Research*, 6(6), e35807.
- [9] Kadirvelu, B., Goel, A., & Choudhury, T. (2025). Multimodal machine learning for mental health prediction: A review. *ACM Computing Surveys*, 57(3), 1–28.
- [10] Angelov, P. P., Soares, E. A., Jiang, R., Arnold, N. I., & Atkinson, P. M. (2021). Explainable artificial intelligence: An analytical review. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 11(5), e1424.