

YOUNG'S MODULUS OF CALCIUM-ALUMINO-SILICATE GLASSES: INSIGHT FROM MACHINE LEARNING

MOUNA Sbai Idrissi¹, Ahmed El Hamdaoui¹, Tarik Chafiq², Said Ouaskit¹

¹Département de physique, Laboratoire de Physique de la Matière Condensée (LPMC), Faculté des Sciences Ben M'Sik, Université Hassan II de Casablanca, BP 7955, Sidi Othman, Casablanca, MAROC, email : mounasbaiid@gmail.com

²Département de géologie, Laboratoire de Physique de la Matière Condensée (LPMC), Faculté des Sciences Ben M'Sik, Université Hassan II de Casablanca, BP 7955, Sidi Othman, Casablanca, MAROC

Abstract: Modern technologies require the development of new materials with exceptional properties. Machine Learning (ML) and Deep Learning (DL) techniques have become important tools for discovering new materials and predicting the properties of specific materials, such as glasses. In this paper, we used ML and DL techniques to predict the Young's modulus E of Calcium-Alumino-Silicate (CAS) glasses based on their chemical composition. We evaluated four different algorithms, including Polynomial Regression (PR), Random Forest (RF), K-Nearest Neighbors (KNN), and Multi-Layer Perceptron Regressor (MLPRegressor). We found that the PR algorithm provides excellent predictions without Cross-Validation (CV), while the MLPRegressor yields the best performance when CV is implemented.

Key words: Calcium-Alumino-Silicate glasses, Young's modulus, Machine Learning (ML), Deep Learning (DL), Cross-Validation (CV).

1. INTRODUCTION

Improving the mechanical properties of glasses, particularly their stiffness (e.g., Young's modulus E), is highly important. Glasses with high stiffness exhibit enhanced resistance to shocks and loads, rendering them safer for use in various applications such as laptop screens, mobile phones, and protective eyewear [1]. Moreover, they find wide applications in industries requiring high optical precision, including scientific instruments [2], construction, and electronic equipment [3]. Additionally, these materials play crucial roles in the immobilization of nuclear waste and aerospace technologies. [4], [5], [6]

Enhancing the stiffness of glasses is one of the most formidable tasks in the realm of physics. [3] Overcoming this challenge necessitates the discovery of new glass compositions with better-suited mechanical properties [7], [8].

However, this endeavour is exceedingly difficult due to the lack of an ordered crystalline structure and fixed stoichiometry in amorphous materials, compounded by their out-of-thermodynamic-equilibrium nature [9], which has a significant impact on the properties of glasses.

In addition to traditional experimental and simulation techniques, deep learning (DL) and machine learning (ML) have emerged as a popular method for predicting the mechanical properties of glasses [10], [11], [12]. Particularly when compared to Molecular Dynamics (MD) simulations, ML and DL methods offer

accurate results albeit with significantly reduced computational time [13]. While MD simulations may require weeks or even months on supercomputers to simulate the mechanical properties of numerous glasses with varied chemical compositions, ML and DL techniques can accomplish the same task within seconds or minutes. [13].

Notably, both molecular dynamics simulations and artificial intelligence simulations have their own advantages and disadvantages, and they can be considered complementary methods. In this study, we apply ML and DL techniques to predict the Young's modulus of the Calcium-Alumino-Silicate (CAS) glass system, with and without cross-validation, and compare the performance of the applied models. The CAS system stands as one of the most extensively studied silicate glass-forming systems, both experimentally and through simulation, owing to its unique optical and mechanical properties, which render it industrially relevant. [14], [15], [16].

This paper is structured into three sections.

The first section focuses on the methodology of the learning algorithms applied to our dataset for predicting the Young's modulus.

The second section presents the results obtained and their interpretations for each model.

Lastly, the concluding section summarizes the findings.

2. METHODOLOGY

2.1 Data Set:

The data set used for prediction includes Young's modulus values for a set of 231 different chemical compositions of $(\text{CaO})_x(\text{Al}_2\text{O}_3)_y(\text{SiO}_2)_{1-x-y}$ glassy systems, with $0 \leq x \leq 1$, $0 \leq y \leq 1$ incremented by 0.05 for both parameters, as depicted in the ternary diagram shown in figure 1. To explain the input variables (or features) of our data set, one should first explain the chemical composition of the CAS glasses, which are composed of Alumino-Silicates with Calcium cations to form a ternary system $\text{CaO}/\text{Al}_2\text{O}_3/\text{SiO}_2$. Therefore, our data set contains the concentrations (% mol) for the three chemical compositions, with an increment of 5%. The only output variable (or target) is Young's modulus E in GPa. The 231 values of E were obtained from high-throughput molecular dynamics simulations (DM) study of CAS glasses by Yang and al[17]. The polymeric structure of these glasses is based on the presence of an aluminate and silicate tetrahedral network that shares corners.

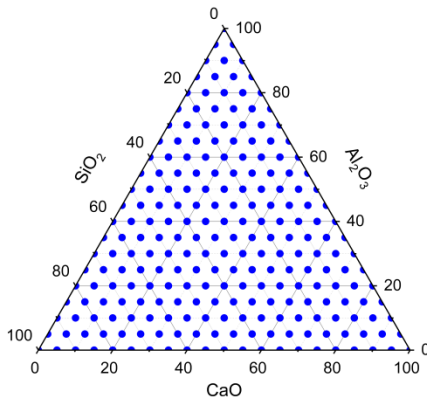


Figure 1 Ternary diagram of the chemical composition of CAS glasses obtained from high-throughput molecular dynamics simulations by Yang and al [17]: Visualization of the percentage concentration of silica, alumina, and calcium oxide elements.

2.2 Learning Algorithms:

We implemented four supervised learning algorithms on our data set: K-Nearest Neighbors (KNN), Random Forest (RF), Polynomial Regression (PR), and Multi-Layer Perceptron (MLP).

The data was divided into a training set (75%) and a test set (25%) with a random splitting. The performance of the models was evaluated based on this initial split. Following the initial evaluation, we employed K-fold Cross-Validation to assess the effect of cross-validation

on each algorithm's performance. this technique further divide the training set into sub-training and validation sets[18], [19]. In this process, the training data was divided into four equal-sized blocks, and during each iteration, one block was used as the validation set, while the other three blocks were used for training[18], [19], [20]. The performance of the model was then evaluated on the current validation set and the overall performance was determined by combining the results from each iteration.

This approach allowed us to observe how the inclusion of cross-validation impacted the behavior and performance of each algorithm.

It's important to note that the initial study involved algorithms evaluated without cross-validation, with performance assessment solely based on the initial training and test split.

2.3 Data Evaluation:

In this work, we evaluate the performance of four different learning algorithms: (KNN), (RF), (PR), and (MLP). Our goal is to understand each model's behavior and performance and to find the best hyperparameters for each algorithm, this involves finding the number of nearest neighbors, the number of trees, the polynomial degree, and the number of neurons for KNN, RF, PR, and MLP respectively. To evaluate the performance of the different models, we use two metrics: the Root Mean Squared Error (RMSE) (Eq.1), and the coefficient of determination R^2 which is the squared correlation coefficient (Eq.2). The model that offers the best performance is the one with the lowest test RMSE and an R^2 close to 1 [21].

$$\text{RMSE} = \sqrt{\frac{1}{n} \sum_{i=1}^n (Y_i - \hat{Y}_i)^2} \quad (1)$$

Correlation coefficient :

$$R = \frac{\sum_{i=1}^n (Y_i - \bar{Y})(\hat{Y}_i - \bar{\hat{Y}})}{\sqrt{\sum_{i=1}^n (Y_i - \bar{Y})^2 \sum_{i=1}^n (\hat{Y}_i - \bar{\hat{Y}})^2}} \quad (2)$$

Coefficient of determination = R^2

Y_i , $\hat{Y}_i$, n , $\bar{Y}$, $\bar{\hat{Y}}$, are the measured (real) value of observation i , the predicted value of observation i by the model, the number of observations, the mean of the measured (real) value of observation i , and the mean of the predicted value of observation i by the model, respectively.

3. RESULTS

3.1 K-Nearest Neighbors:

The KNN model uses the number of nearest neighbors (k) as an important hyperparameter [22]. Fig.2(a) shows the RMSE values for the test and training

sets as a function of k , with and without CV. For low values of k , we observe that the RMSE of test set are higher than that of training set, which indicate that the model is overfitting to the training data and not generalizing well to unseen data [23]. By increasing the value of k , the RMSE of the training set increases and the RMSE of the test set decreases until reaching a value of k equal to 4. Beyond 4, we observed that the RMSE of both the training and test sets increased, indicating that the model is unable to fit the training data well enough to offer good performance on the test data, which is known as underfitting. In this model, the complexity varies depending on several factors, including the number of neighbors (k) considered for predicting a point, the dimensionality of the space in which the data resides, and the total number of points in the dataset. Regarding the number of neighbors (k), the complexity of the KNN model generally increases with k . This is because in order to predict the class of a point, the model needs to find the k nearest points in the data set, which can be computationally expensive if k is large. However, a value of k that is too small can also lead to insufficient model performance [22], [23], [24], [25]. Indeed, we found that k equal to 4 provides a minimum for the test RMSE, which is 3.36 (GPa). Therefore, k equal to 4 is the optimal value for this model. We now focus on evaluating the accuracy of the KNN model using the determination coefficient R^2 for the test set, with and without CV. Figure 2(b) shows that for $k=4$, the highest R^2 scores without cross-validation and with cross-validation are respectively 0.96 and 0.94.

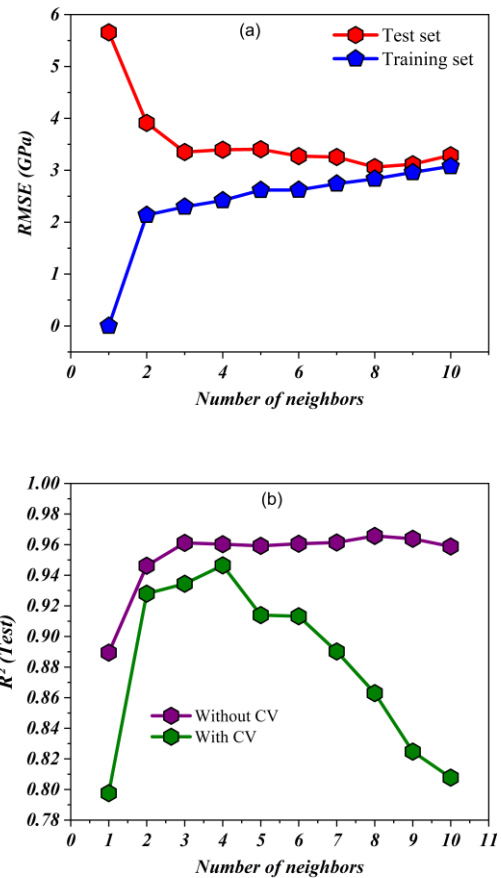


Figure 2

(a) RMSE of the KNN model as a function of the number of nearest neighbors as obtained for the training and test set;

(b) The influence of the number of nearest neighbors on the performance of the KNN model: Comparison of the determination coefficients obtained with and without CV.

3.2 Random Forest:

In this model, we focused on the effect of the number of trees in a random forest on performance. The results illustrated in Figure 3(a) demonstrate that the RMSE for the test and training sets shows a small variation with the number of trees, with and without CV. No overfitting or underfitting is observed, which is due to the simplicity of this model. We find 100 trees as the optimal hyperparameter that provides a minimum RMSE [20], [26] of 2.904 GPa for the test data. When evaluating the accuracy of the (RF) model using the R^2 determination coefficient for the test set, Figure 3(b) shows that the highest scores are obtained when the number of trees used is equal to 100. Using CV, the score obtained is 0.933, while in the absence of CV, the score obtained is 0.970. These results show that the use

of CV has an impact on the model accuracy, but the number of trees used remains an important factor in achieving the best performance.

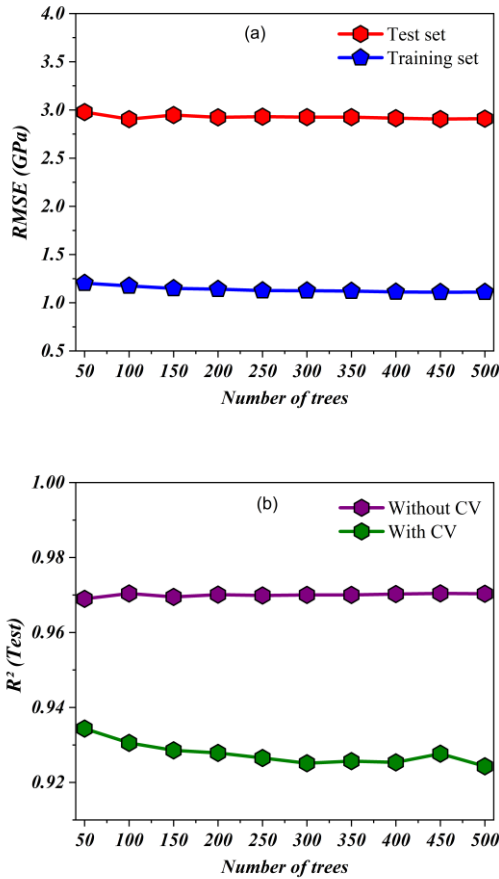


Figure 3

(a) RMSE of the RF model as a function of the number of trees as obtained for the training and test set;

(b) Influence of the number of trees on the performance of the RF model: Comparison of the determination coefficients obtained with and without CV.

3.3 Polynomial Regression:

We used the PR to train our data, selecting polynomial degrees ranging from 1 to 10. Figure 4(a) depicts the RMSE for both training and test data as a function of polynomial degrees. The RMSE values for training data are clearly lower than those for test data as shown in Figure 4(a). For low polynomial degrees, the high RMSE values for both sets indicate that the model is too simple to accurately represent the relationship between Young's modulus and the chemical composition of CAS glasses, leading to underfitting. As the polynomial degree increases, the RMSE values decrease

until reaching a minimum value of 2.41 GPa for the test data at degree 5. Beyond this point, the test RMSE begins to increase again, which means the model has overfit. Therefore, the optimal polynomial degree is 5, which offers the best performance both with and without CV Figure 4(b).

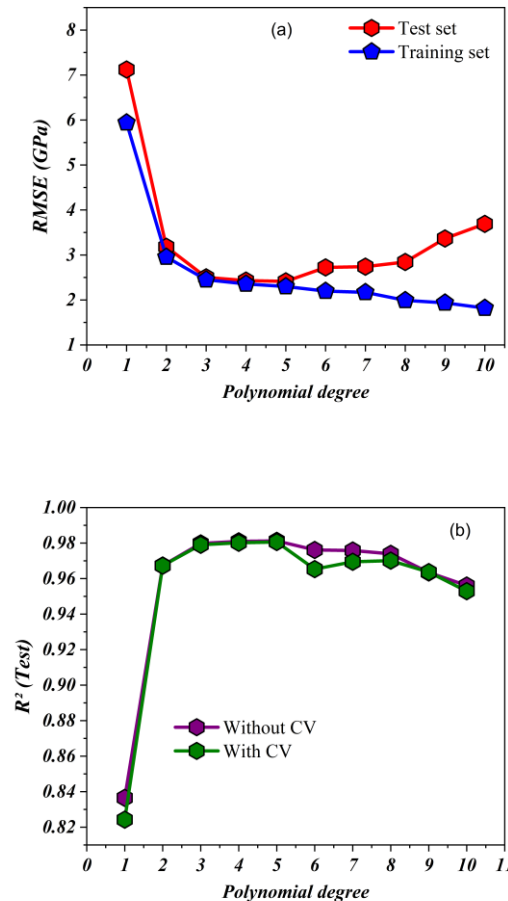


Figure 4

(a) RMSE of the PR model as a function of the polynomial degree for training and test sets;

(b) Influence of the polynomial degree on the performance of the PR model: Comparison of determination coefficients obtained with and without CV.

3.4 Multilayer Perceptron:

This DL model based on neurons [27], [28], [29], [30] consists of a single hidden layer with a number of neurons ranging from 50 to 500, incremented by 50, and a single output neuron. The activation function chosen for the hidden layer is the "relu" function that returns " $f(x) = \max(0, x)$ ", the maximum between 0 and x [31]. Figure 5(a) shows the performance of the MLP model in terms of RMSE for the training and test sets as a

function of the number of neurons, with and without CV. It is clear that the MLP model does not show notable overfitting. When the number of neurons is 50, the RMSE for the test data is at a minimum, which means that 50 neurons is the optimal hyperparameter for this model. When evaluating the performance of the MLP model using the R^2 determination coefficient for the test set, Figure 5(b) shows that the highest scores are obtained with 50 and 150 neurons respectively without and with CV. The use of CV increased the model's accuracy score from 0.975 to 0.987. These results indicate that the number of neurons used and the use of CV have a significant impact on the accuracy of the MLP model.

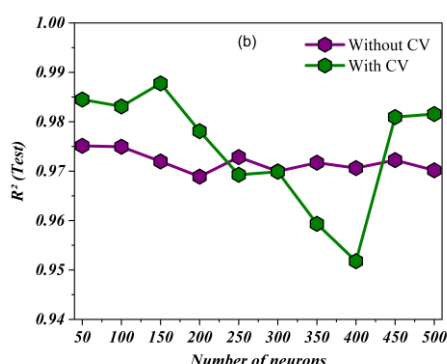


Figure 5

(a) RMSE of the MLP model as a function of the number of neurons for training and test sets;
(b) The influence of the number of neurons on the performance of the MLP model: Comparison of determination coefficients obtained with and without CV.

4. DISCUSSION

In this section, we will discuss the results obtained by different DL and ML algorithms with and without CV. Table 1. presents the R^2 determination coefficient values obtained for each algorithm on the test set, which measures the algorithm's ability to predict unknown data and thus the accuracy of the model. We found that the KNN-algorithm offers an accuracy of 96% without CV and 94% with CV. The RF-algorithm of Trees model gives an accuracy of 97% without CV and 93% with CV. The MLP and PR algorithms have a similar accuracy, although the PR-algorithm has slightly higher accuracy at 98.12% without CV. Finally, we observe that the MLP-algorithm provides the highest accuracy with a CV of 98.77%.

It is important to note that ML algorithms are generally easier to understand and interpret than deep neural networks, but the latter tend to give more accurate results in many cases, such as image recognition[28]. It is therefore important to choose the right type of algorithm depending on the precision requirements. The

KNN, RF, and PR machine learning models have better performance without CV than with it. This may be due to the small size of the used data set (231 samples) [32], [33] which can reduce accuracy when CV is used. It is also possible that the CV for training data is more representative of the test data set, which can result in better performance. However, for some algorithms such as RF, the overfitting is not observed, which means it can work without CV [32].

It is important to note that CV is a technique used to evaluate the performance of models using sub-training data that has not been used for training, but it is not always necessary or beneficial for all models and data sets. For DL models such as MLP, the performance is usually better with CV due to the complexity of these models[34].

Figures.6(a)-(b) show the predicted young's modulus (E) values by PR without cross-validation (with a degree of 5) and by MLP with cross-validation (with 150 neurons) in comparison with the simulated young's modulus values by high-throughput molecular dynamics simulations in (GPa), respectively. Therefore, these models are the most effective and can ameliorate the prediction of the relationship between the chemical composition and young's modulus of the CAS glass system.

Table 1. Comparison of performance of different learning algorithms: Coefficient of determination obtained without and with CV for the entire test data for each model

Algorithmes d'apprentissage	R^2 (test)	
	Sans CV	Avec CV
KNN	0.960	0.946
RF	0.970	0.930
PR	0.981	0.980
MLP	0.975	0.984

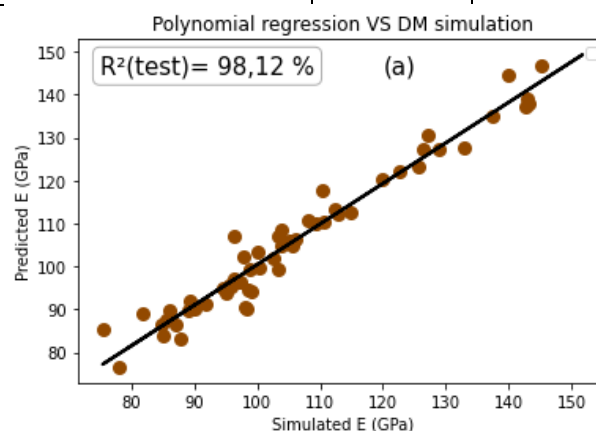


Figure 6 (a) Comparison between the predicted Young's modulus values by PR (degree = 5) without cross-validation and those calculated by molecular dynamics simulations

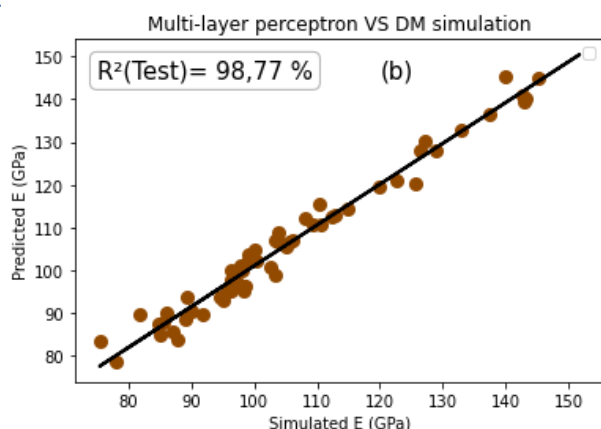


Figure 6 (b) Comparison between the predicted young's modulus values by MLP with cross-validation (number of neurons = 150) and those calculated by molecular dynamics simulations.

5. CONCLUSION

Overall, this paper examines how artificial intelligence approaches can be used to discover and predict new mechanical properties of glasses (Young's modulus) without relying on computer simulation techniques based on numerical calculations. The learning algorithms are flexible and easy to use.

Upon application to our dataset, the MLP algorithm with cross-validation exhibits the highest performance. Conversely, the PR algorithm achieves optimal results without cross-validation, specifically with a polynomial degree of 5. This disparity in performance highlights an interesting observation regarding the impact of cross-validation on machine learning algorithms compared to deep learning algorithms. The simplicity of the machine learning algorithms appears to render them less sensitive to the effects of cross-validation, whereas deep learning algorithms may benefit more significantly from such validation techniques. This distinction underscores the nuanced relationship between model complexity, algorithmic simplicity, and the efficacy of cross-validation methodologies in optimizing predictive performance.

However, further applications of ML and DL algorithms are required to examine different types of silicate glasses (binary and ternary) in order to reach a general conclusion on the use of these methods in predicting the properties of glassy materials.

6. REFERENCES

[1] S. Bishnoi *et al.*, "Predicting Young's modulus of oxide glasses with sparse datasets using machine learning," *Journal of Non-Crystalline Solids*, vol. 524, p. 119643, Nov. 2019, DOI 10.1016/j.jnoncrsol.2019.119643.

[2] "Bulk elastic properties, hardness and fatigue of calcium aluminosilicate glasses in the intermediate-silica range - ScienceDirect." Accessed: Mar. 05, 2024. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0022309315302829>

[3] J. C. Mauro and E. D. Zanotto, "Two Centuries of Glass Research: Historical Trends, Current Status, and Grand Challenges for the Future," *International Journal of Applied Glass Science*, vol. 5, no. 3, pp. 313–327, 2014, doi: 10.1111/ijag.12087.

[4] "Glasses-for-Nuclear-Waste-Immobilization.pdf." Accessed: Mar. 05, 2024. [Online]. Available: https://www.researchgate.net/profile/Michael-Ojovan/publication/267700284_Glasses_for_Nuclear_Waste_Immobilization/links/5465a9240cf2f5eb17ff42de/Glasses-for-Nuclear-Waste-Immobilization.pdf

[5] C. I. Merzbacher, B. L. Sherriff, J. S. Hartman, and W. B. White, "A high-resolution ^{29}Si and ^{27}Al NMR study of alkaline earth aluminosilicate glasses," *Journal of Non-Crystalline Solids*, vol. 124, no. 2, pp. 194–206, Oct. 1990, doi: 10.1016/0022-3093(90)90263-L.

[6] M. Toozandehjani, N. Kamarudin, Z. Dashtizadeh, E. Y. Lim, A. Gomes, and C. Gomes, "Conventional and Advanced Composites in Aerospace Industry: Technologies Revisited," *American Journal of Aerospace Engineering*, vol. 5, pp. 9–15, Feb. 2018, doi: 10.11648/j.ajae.20180501.12

[7] "Decoding the glass genome - ScienceDirect." Accessed: Mar. 05, 2024. [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S1359028617301249>

[8] "Machine learning predictions of Knoop hardness in lithium disilicate glass-ceramics, Wilkinson, 2023, *Journal of the American Ceramic Society*, Wiley Online Library. Accessed: Mar. 05, 2024. [Online]. Available: <https://ceramics.onlinelibrary.wiley.com/doi/abs/10.1111/jace.19016>

[9] A. K. Varshneya, *Fundamentals of Inorganic Glasses*. Elsevier, 2013.

[10] A. W. Abboud, D. P. Guillen, and B. A. Christensen, "Prediction of melter cold-cap topology from plenum temperatures with computational fluid dynamics and machine learning," *International Journal of Ceramic Engineering & Science*, vol. 4, no. 4, pp. 257–269, 2022, doi: 10.1002/ces2.10134.

[11] M. A. Kraus and M. Drass, "Artificial intelligence for structural glass engineering applications —



- overview, case studies and future potentials,” *Glass Struct Eng*, vol. 5, no. 3, pp. 247–285, Nov. 2020, doi: 10.1007/s40940-020-00132-8.
- [12] M. Drass, H. Berthold, M. A. Kraus, and S. Müller-Braun, 2021. “Semantic segmentation with deep learning: detection of cracks at the cut edge of glass,” *Glass Struct Eng*, vol. 6, no. 1, pp. 21–37, doi: 10.1007/s40940-020-00133-7.
- [13] J. Singh and S. Singh, 2022, “A review on Machine learning aspect in physics and mechanics of glasses,” *Materials Science and Engineering: B*, vol. 284, p. 115858, Oct. 2022, DOI: 10.1016/j.mseb.2022.115858.
- [14] “Formation and Properties of Calcium Aluminosilicate Glasses”, SHELBY, 1985, *Journal of the American Ceramic Society*, Wiley Online Library. Accessed: Mar. 05, 2024. [Online]. Available: <https://ceramics.onlinelibrary.wiley.com/doi/abs/10.1111/j.1151-2916.1985.tb09656.x>
- [15] M. E. Lines, J. B. MacChesney, K. B. Lyons, A. J. Bruce, A. E. Miller, and K. Nassau, 1989, “Calcium aluminate glasses as potential ultralow-loss optical materials at 1.5–1.9 μm ,” *Journal of Non-Crystalline Solids*, vol. 107, no. 2, pp. 251–260, Jan. 1989, doi: 10.1016/0022-3093(89)90470-5.
- [16] F. T. Wallenberger and S. D. Brown, “High-modulus glass fibers for new transportation and infrastructure composites and new infrared uses,” *Composites Science and Technology*, vol. 51, no. 2, pp. 243–263, Jan. 1994, Doi: 10.1016/0266-3538(94)90194-5.
- [17] K. Yang et al., 2019, “Predicting the Young’s Modulus of Silicate Glasses using High-Throughput Molecular Dynamics Simulations and Machine Learning,” *Sci Rep*, vol. 9, no. 1, p. 8739, Jun. 2019, Doi: 10.1038/s41598-019-45344-3.
- [18] R. Kohavi, 2001, “A Study of Cross-Validation and Bootstrap for Accuracy Estimation and Model Selection,” vol. 14.
- [19] Z. Nematzadeh, *Comparative Studies on Breast Cancer Classifications with K-Fold Cross Validations Using Machine Learning Techniques*. 2015.
- [20] G. Battineni, G. G. Sagaro, C. Nalini, F. Amenta, and S. K. Tayebati, 2019, “Comparative machine-learning approach: A follow-up study on type 2 diabetes predictions by cross-validation methods,” *Machines*, vol. 7, no. 4, p. 74.
- [21] T. Oey, S. Jones, J. W. Bullard, and G. Sant, 2020. “Machine learning can predict setting behavior and strength evolution of hydrating cement systems,” *J Am Ceram Soc*, vol. 103, no. 1, pp. 480–490, Doi: 10.1111/jace.16706.
- [22] Z. Zhang, 2024, “Introduction to machine learning: k-nearest neighbors,” *Annals of translational medicine*, vol. 4, no. 11, 2016, Accessed: Mar. 06, 2024. [Online]. Available: <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC4916348/>
- [23] “Overfitting and Underfitting Analysis for Deep Learning Based End-to-end Communication Systems” *IEEE Conference Publication | IEEE Xplore*. Accessed: Mar. 06, 2024. [Online]. Available: <https://ieeexplore.ieee.org/abstract/document/8927876>
- [24] S. K. Ahmmad, N. Jabeen, S. T. U. Ahmed, S. A. Ahmed, and S. Rahman, 2021, “Artificial intelligence density model for oxide glasses,” *Ceramics international*, vol. 47, no. 6, pp. 7946–7956.
- [25] S. Huang, M. Huang, and Y. Lyu, 2020, “An Improved KNN-Based Slope Stability Prediction Model,” *Advances in Civil Engineering*, vol. 2020, p. e8894109, Jul. 2020, doi: 10.1155/2020/8894109.
- [26] L. Yang, K. Jia, S. Liang, J. Liu, and X. Wang, 2016, “Comparison of four machine learning methods for generating the GLASS fractional vegetation cover product from MODIS data,” *Remote Sensing*, vol. 8, no. 8, p. 682, 2016.
- [27] A. Nguyen, K. Pham, D. Ngo, T. Ngo, and L. Pham, 2021, “An analysis of state-of-the-art activation functions for supervised deep neural network,” in *2021 International Conference on System Science and Engineering (ICSSE)*, IEEE, pp. 215–220. Accessed: Mar. 06, 2024. [Online]. Available: <https://ieeexplore.ieee.org/abstract/document/9538437/>
- [28] D. Jakhar and I. Kaur, 2020, “Artificial intelligence, machine learning and deep learning: definitions and differences,” *Clinical and experimental dermatology*, vol. 45, no. 1, pp. 131–132.
- [29] S. Abirami and P. Chitra, 2020, “Energy-efficient edge based real-time healthcare support system,” in *Advances in computers*, vol. 117, Elsevier, pp. 339–368. Accessed: Mar. 06, 2024. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0065245819300506>
- [30] I.-C. Yeh, 1998 “Modeling of strength of high-performance concrete using artificial neural networks,”



Cement and Concrete research, vol. 28, no. 12, pp. 1797–1808.

[31] D. Roche, 2024, “*Deep Learning et apprentissage par renforcement pour la conception d’une Intelligence Artificielle pour le jeu Yokai No Mori*,” Ce travail est publié sous licence Creative Common, Accessed: Mar. 06, 2024. [Online]. Available: <https://www.epi.asso.fr/revue/articles/a1905d.pdf>.

[32] H. Blockeel and J. Struyf, 2002, “*Efficient algorithms for decision tree cross-validation*,” Journal of Machine Learning Research, vol. 3, no. Dec, pp. 621–650.

[33] G. Varoquaux, 2018, “*Cross-validation failure: Small sample sizes lead to large error bars*,” Neuroimage, vol. 180, pp. 68–77.

[34] Jason Brownlee , 2024, “*Better Deep Learning: Train Faster, Reduce Overfitting, and Make Better* , Google Livres. Accessed: Mar. 06, 2024. [Online]. Available: https://books.google.co.ma/books?hl=fr&lr=&id=T1-nDwAAQBAJ&oi=fnd&pg=PP1&dq=Better+deep+learning:+train+faster,+reduce+overfitting,+and+make+better+predictions&ots=tENTfoj0FX&sig=_3tyCki2fS3E9FBOCYvII7M8Or0&redir_esc=y#v=onepage&q=Better%20deep%20learning%3A%20train%20faster%2C%20reduce%20overfitting%2C%20and%20make%20better%20predictions&f=false