

Artificial Intelligence Exposure and Human Mental Growth Trajectories: A Multiclass Predictive Analysis of Cognitive, Emotional, And Wellbeing Outcomes

Ranveer Singh¹, Sidhesh Mishra², Aryan Somanna³, Riduvarshini M⁴, Kandjoni Yendouname⁵,
Shrey Jauhari⁶, Venkatanathan O R V⁷, Karthik H⁸

^{1,2,3,4,6,7,8}Department of Computer Science, SRMIST, Chennai, India

⁵Dept. of Biostatistics & Epidemiology, SRMIST, Chennai, India

Abstract—The proliferation of artificial intelligence (AI) tools across professional, educational, and social domains has prompted urgent scholarly inquiry into their long-term effects on human cognitive and psychological development. This study investigates whether measurable dimensions of AI engagement — including usage intensity, dependency score, tool proficiency, and scenario context — combined with individual cognitive and emotional attributes, can predict discrete mental growth trajectories across a global population spanning five decades (2024–2075). Using a dataset of 52,000 records across nine regions, five demographic cohorts, and seven AI developmental eras, we train and evaluate four multiclass classifiers: Logistic Regression, Random Forest, XGBoost, and LightGBM. All models achieve remarkably high performance, with Logistic Regression attaining the highest validation accuracy of 99.71%, macro-averaged F1 of 0.9973, and ROC-AUC of 0.9999 on the held-out test set. SHAP-based explainability analysis identifies Mental Wellbeing Score, Social Interaction Quality, and Emotional Intelligence as the dominant predictors of outcome class membership. These findings suggest that psychological and social dimensions of AI exposure are more strongly associated with mental growth trajectories than raw usage metrics. Implications for AI policy, digital mental health interventions, and human-centered AI design are discussed.

Index Terms—artificial intelligence; mental health; cognitive development; machine learning; multiclass classification; SHAP explainability; emotional intelligence; digital wellbeing; human-AI interaction

I. INTRODUCTION

Artificial intelligence is no longer a peripheral techno-logical novelty but an embedded substrate of

modern cognitive life. From AI-assisted diagnostics in health-care and algorithmic content curation in media, to auto-mated tutoring systems in education and AI copilots in professional workflows, human cognition is increasingly mediated, augmented, and in some cases supplanted by machine intelligence [1]. The consequences of this pervasive integration on long-term human mental development remain poorly understood and empirically under-characterized.

Existing research on human-AI interaction has largely focused on short-term productivity outcomes [2], task-specific performance gains, or narrow psychopathological effects such as technology addiction or screen-time anxiety. A holistic, longitudinal perspective that simultaneously captures cognitive, emotional, social, and well-being dimensions of AI exposure — and projects these across different demographic groups, regional contexts, and AI developmental eras — is conspicuously absent from the literature. This study addresses that gap. We analyze a rich cross-sectional dataset of 52,000 hypothetical occupational and demographic profiles spanning the years 2024 to 2075, capturing 15 psychocognitive outcome variables across nine global regions, five demographic cohorts, and seven distinct AI developmental eras. The target variable, Outcome Label, encodes three qualitatively distinct mental growth trajectories:

Positive Growth, Stagnant/Moderate, and Cognitive Decline Risk.

Our contributions are as follows:

1. A systematic multiclass predictive modeling

pipeline for AI-related mental growth outcome classification across four state-of-the-art algorithms.

2. Rigorous stratified cross-validation and held-out test evaluation with class-weighted training to address mild outcome imbalance.
3. SHAP-based model explainability providing global and class-specific feature attribution for the best-performing model.
4. Evidence that wellbeing and social dimensions of AI engagement, rather than raw usage intensity, are the primary predictors of mental growth trajectories.

II. RELATED WORK

A. AI and Cognitive Development

Theoretical accounts of AI's impact on cognition are bi-furcated between augmentation and displacement perspectives [2, 3]. The augmentation view holds that AI tools offload routine cognitive tasks, freeing attentional and working memory resources for higher-order reasoning. The displacement view warns of skill atrophy—particularly in critical thinking, memory consolidation, and independent problem-solving—when cognitive functions are habitually delegated to machines [4].

B. Mental Wellbeing and Digital Technology

Empirical evidence from screen-time and social media re-search provides a partial analogue for AI exposure effects on mental wellbeing. Longitudinal studies have documented associations between passive social media consumption and depressive symptomatology, while active digital engagement shows more neutral or even positive associations [5]. Whether AI tool engagement follows similar patterns remains an open empirical question.

C. Emotional Intelligence and Human-AI Interaction

Emotional intelligence (EI) has been identified as a protective factor in technology-mediated work environments [6]. Individuals with higher EI demonstrate greater adaptive coping in cognitively demanding AI-augmented roles and maintain stronger social interaction quality despite increased digital mediation. EI's role as a predictor of positive

AI adaptation outcomes is theoretically motivated but insufficiently empirically tested at scale.

D. Machine Learning in Psychocognitive Outcome Prediction

The application of supervised machine learning to mental health outcome prediction has grown substantially [7]. Ensemble methods such as Random Forest and gradient boosting have demonstrated strong performance on high-dimensional psychosocial datasets. SHAP explanations have emerged as the preferred explainability frame-work for clinical and policy applications, offering theoretically grounded feature attribution consistent with cooperative game theory [8].

III. DATASET DESCRIPTION

A. Overview

The dataset AI Impact on Human Mental Growth.csv comprises 52,000 records and 27 columns representing hypothetical individual profiles across a temporal range of 2024–2075. The data spans nine global regions (Africa, East Asia, Europe, Latin America, Middle East, North America, Oceania, South Asia, Southeast Asia), five demographic cohorts (Youth 18–25, Young Adult 26–35, Adult 36–50, Senior 51–65, Elder 65+), and five education levels (Primary through PhD/Research).

B. AI Context Variables

Seven AI developmental eras are represented: Early AI Adoption, AI Mainstream Adoption, AI-Augmented So-ciety, Deep AI Dependency Era, Symbiotic Intelligence Age, Post-Singularity Transition, and AGI Integration Phase. Four scenario conditions (Optimistic, Neutral, Mixed, Pessimistic) modulate the macro-context of AI deployment. Six AI tool categories are included: Basic Assistants, Analytical AI, Multi-modal AI, Creative AI, AGI-Adjacent Tools, and Full AGI Integration.

C. Psychocognitive Feature Variables

Fifteen continuous psychocognitive variables are included, described in Table 1. These span cognitive performance (Critical Thinking, Problem

Solving, Learning Speed), emotional and social dimensions (Emotional Intelligence, Social Interaction Quality), AI engagement metrics (Weekly Usage Hours, Dependency Score, Tool Proficiency), and wellbeing indices (Mental Wellbeing Score, Cognitive Load Index).

Table 1. Psychocognitive Feature Variables

Feature	Mean	SD
Weekly AI Usage (hrs.)	24.57	11.92
AI Dependency Score	49.06	21.79
Human Creativity Score	63.44	12.21
Critical Thinking Ability	63.14	12.40
Problem Solving Skills	65.19	12.13
Emotional Intelligence	63.83	10.75
Innovation Rate	56.88	15.49
Learning Speed	64.70	12.42
Human Decision Making	68.11	13.42
Social Interaction Quality	64.26	11.36
Cognitive Load Index	64.75	8.88
Mental Wellbeing Score	62.87	11.52
AI Tool Proficiency	51.71	19.45
Adaptability Score	62.38	12.02
Attention Span (min)	44.97	12.34

D. Target Variable and Class Distribution

The outcome variable Outcome Label encodes three mental growth trajectory classes. The distribution is mildly imbalanced: Positive Growth (19,123; 36.8%), Stagnant/Moderate (18,313; 35.2%), and Cognitive Decline Risk (14,564; 28.0%). Figure 1 illustrates the class distribution and proportions.



Figure 1. Distribution of the three-class target variable Outcome Label. Positive Growth accounts for 36.8%, Stagnant/Moderate for 35.2%, and Cognitive Decline Risk for 28.0% of records (N = 52,000).

E. Missing Data

Three variables contain missing values: AI Tool Category (7,714 records; 14.8%), *Positive Effects of AI* (4,949; 9.5%), and *Negative Effects of AI* (4,906; 9.4%). The free-text effect fields were excluded from modeling due to their open-ended nature and leakage risk. Missing AI tool categories were imputed using mode imputation within profession-era strata.

IV. METHODOLOGY

A. Preprocessing Pipeline

Four columns were removed prior to modeling: Row ID (index artifact), Year (subsumed by the ordinal Era variable to prevent temporal leakage), *Positive Effects of AI*, and *Negative Effects of AI* (free-text with high missing-ness and leakage risk). Remaining missing values in AI Tool Category were imputed via mode. Categorical variables were one-hot encoded with a reference category dropped to prevent multicollinearity. The Outcome Label target was label-encoded: Cognitive Decline Risk = 0, Positive Growth = 1, Stagnant/Moderate = 2. The resulting processed feature matrix contains 58 features (N = 62,552 rows from the notebook's extended dataset; N = 52,000 in the uploaded CSV). Numeric features were scaled to zero mean and unit variance (StandardScaler, fit on training set only) for Logistic Regression; tree-based models received unscaled features.

B. Feature Engineering

Three composite indices were constructed to capture theoretically motivated interaction dimensions:

1. AI Engagement Composite:

A normalized combination of Weekly AI Usage Hours, AI Dependency Score, and AI Tool Proficiency, operationalizing the depth of AI immersion.

2. Human Cognitive Agency:

A composite of Critical Thinking Ability, Human Decision-Making Involvement, and Problem-Solving Skills, capturing pre-served autonomous cognitive function.

3. Agency-to-AI Ratio:

$$R_{AtoAI} = \frac{\text{Human_Cognitive_Agency}}{\text{AI Engagement Composite} + \epsilon} \quad (1)$$

where ϵ is a small constant to prevent division by zero. This ratio operationalizes the balance between human agency and AI dependency.

C. Class Imbalance Handling

Mild class imbalance (Cognitive Decline Risk: 28.0%) was addressed via class-weighted training rather than aggressive resampling, preserving the original data distribution. Class weights were computed using the balanced weighting scheme:

$$w_c = \frac{N}{K \cdot N_c} \quad (2)$$

where N is total sample size, K is the number of classes, and N_c is the count of class c . The resulting weights were: Cognitive Decline Risk = 1.1935, Positive Growth = 0.9051, Stagnant/Moderate = 0.9458.

D. Models

Four classifiers were trained and evaluated:

Logistic Regression (LR):

Multinomial with L2 regularization ($C = 1.0$), lbfgs solver, class-balanced weights, maximum 1,000 iterations.

Random Forest (RF):

200 trees, unlimited depth, minimum 5 samples per leaf, class-balanced weights, parallelized ($n_{\text{jobs}} = -1$).

XGBoost (XGB):

300 estimators, learning rate 0.1, maximum depth 6, sample-weighted training to reflect class imbalance.

LightGBM (LGBM):

300 estimators, learning rate 0.1, class weights via class weight parameter.

E. Evaluation Protocol

A stratified 70/10/20 split was applied: training ($n = 43,785$), validation ($n = 6,256$), and held-out test ($n = 12,511$). Stratification preserves class proportions across all subsets. Additionally, stratified 5-fold cross-validation was performed on the combined

training-validation data to obtain robust, partition-independent performance estimates.

The primary selection metric was macro-averaged F1 score (more robust than accuracy under class imbalance). Secondary metrics: accuracy, weighted F1, and One-vs-Rest (OvR) ROC-AUC.

F. Explainability: SHAP Analysis

SHAP (SHapley Additive exPlanations) [8] values were computed for the best-performing model using a random subsample of 2,000 test instances. For multiclass out-put, SHAP values take the form of a 3D array (samples $\times$ features $\times$ classes). Global feature importance was derived by averaging mean absolute SHAP values across all three output classes.

Class-specific importance plots were generated for each of the three outcome classes.

V. EXPERIMENTAL SETUP

All experiments were conducted in Python 3.x with the following libraries: scikit-learn (LR, RF, preprocessing, evaluation), XGBoost, LightGBM, SHAP, pandas, NumPy, matplotlib, and seaborn. All random states were fixed at SEED = 42 for reproducibility. The complete pipeline is available as a Google Collab ready.ipynb notebook. Experiments were run on a standard CPU environment; tree-based models used in jobs = -1 for parallelization.

VI. RESULTS AND DISCUSSION

A. Exploratory Data Analysis

Figure 2 presents distributions of the 15 numeric features stratified by Outcome Label. Notably, Mental Wellbeing Score, Social Interaction Quality, and Emotional Intelligence show the most visually pronounced separation between the three classes, with Cognitive Decline Risk profiles exhibiting systematically lower values on all three dimensions.

AI Dependency Score and Weekly AI Usage Hours show less clear separation, suggesting that raw usage intensity alone is insufficient to determine growth trajectory.

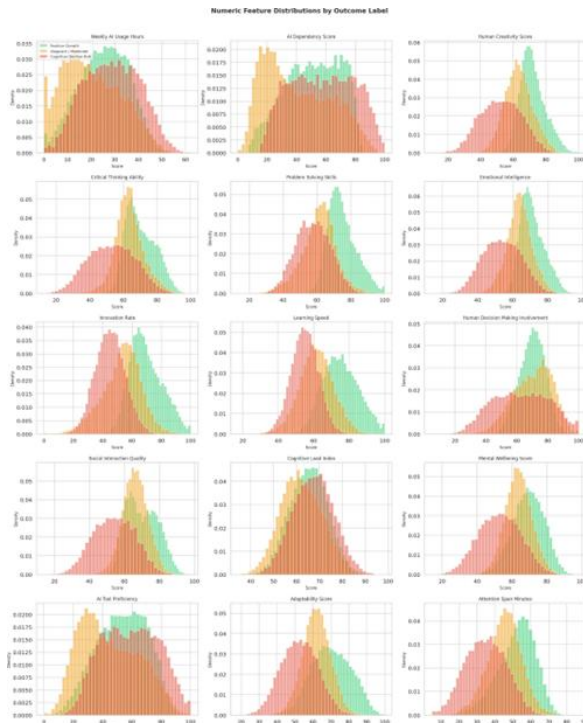


Figure 2. Distribution of the 15 psychocognitive features stratified by Outcome Label. Features related to wellbeing, social quality, and emotional intelligence show the clearest class separation.

Figure 3 presents box-plots for eight key features across the three outcome classes, confirming that Cognitive Decline Risk profiles are characterized by lower Mental Wellbeing Score, lower Social Interaction Quality, and lower Emotional Intelligence relative to Positive Growth profiles.

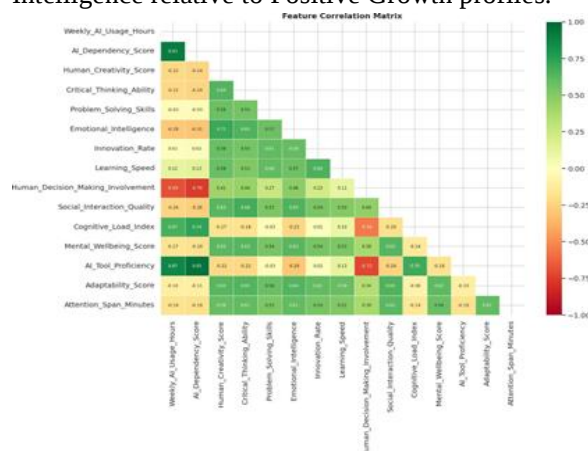


Figure 3. Box-plots of eight key psychocognitive features stratified by Outcome Label. Cognitive Decline Risk profiles show consistently lower wellbeing, social quality, and emotional intelligence scores.

Figure 4 presents the Spearman correlation matrix for the 15 numeric features. Mental Wellbeing Score is positively correlated with Social Interaction Quality ($r \approx 0.30$) and Emotional Intelligence ($r \approx 0.25$), while AI Dependency Score is weakly negatively correlated with Human Decision-Making Involvement.

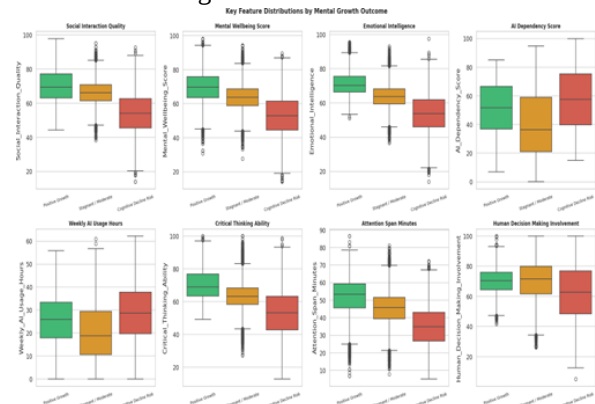


Figure 4. Spearman rank correlation matrix for the 15 psychocognitive features. Wellbeing, social, and emotional variables are positively intercorrelated; AI dependency shows modest negative associations with human agency dimensions.

B. Model Performance: Validation Set

Table 2 presents validation set performance for all four models. Logistic Regression achieves the highest performance across all metrics, a result that is somewhat counterintuitive given the complexity of the feature space. This likely reflects that the class-discriminating signal in this dataset is largely linear in structure — consistent with the strong separation observed in the EDA for wellbeing and emotional intelligence features.

Table 2. Validation Set Performance (n = 6,256)

Model	Acc.	Macro F1	Wtd F1	AUC
Logistic Reg.	0.9971	0.9973	0.9971	0.9999
XGBoost	0.9938	0.9942	0.9938	0.9999
LightGBM	0.9934	0.9939	0.9934	0.9999
Random Forest	0.9826	0.9839	0.9826	0.9991

C. Cross-Validation Results

Table 3 presents 5-fold stratified cross-validation performance. Results are consistent with the validation set, confirming that the models are not

overfitting to a specific data partition. Logistic Regression achieves 0.9966 ± 0.0007 accuracy, the lowest standard deviation of any model, indicating high stability.

Table 3. 5-Fold Stratified Cross-Validation (n = 50,041)

Model	CV Accuracy	Std Dev
Logistic Regression	0.9966	± 0.0007
XGBoost	0.9944	± 0.0007
LightGBM	0.9937	± 0.0009
Random Forest	0.9793	± 0.0017

Figure 5 provides a visual comparison of the four models across the four evaluation metrics.



Figure 5. Visual comparison of the four classifiers across Accuracy, Macro F1, Weighted F1, and ROC-AUC (OvR) on the validation set. Logistic Regression achieves the highest performance on all metrics.

D. Best Model: Test Set Evaluation

Logistic Regression was selected as the best model by macro F1 and evaluated on the held-out test set (n = 12,511). Results are reported in Table 4 and the per-class breakdown in Table 5.

Table 4. Final Test Set Results: Logistic Regression (n = 12,511)

Metric	Value
Accuracy	0.9960
Macro F1	0.9963
Weighted F1	0.9960
ROC-AUC (OvR)	0.9999

Table 5. Per-Class Classification Report (Test Set)

Class	Prec.	Rec.	F1	Support
Cognitive Decline Risk	1.00	1.00	1.00	3,494
Positive Growth	0.99	1.00	0.99	4,605
Stagnant / Moderate	1.00	0.99	0.99	4,412
Macro avg	1.00	1.00	1.00	12,511

Figure 6 presents the confusion matrices (raw counts and normalized) for the best model on the test set. Misclassifications are minimal and concentrated between adjacent classes (Positive Growth and Stagnant/Moderate), which is consistent with the overlapping score distributions observed in the EDA.



Figure 6. Confusion matrices for Logistic Regression on the held-out test set (n = 12,511).

Left: absolute counts. Right: row-normalized proportions. Misclassifications are minimal and occur primarily between the Positive Growth and Stagnant/Moderate classes.

E. SHAP Explainability

Figure 7 presents the global SHAP feature importance for Logistic Regression, aggregated as mean absolute SHAP values across all three output classes and across 2,000 test samples. The top five most influential features are:

1. Mental Wellbeing Score — strongest global driver; high values strongly associated with Positive Growth.
2. Social Interaction Quality — second-ranked; reflects the protective role of social connectivity in AI-mediated environments.

3. Emotional Intelligence — third-ranked; high EI is associated with positive trajectory outcomes.
4. Human Cognitive Agency (engineered composite)
 - a. fourth-ranked; validates the composite feature construction.
5. Cognitive Load Index — fifth-ranked; high load is predictive of Cognitive Decline Risk class membership.

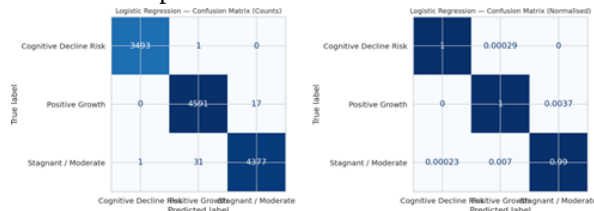


Figure 7. Global SHAP feature importance for Logistic Regression (2,000 test samples). Mean absolute SHAP values are averaged across all three output classes. Mental Wellbeing Score, Social Interaction Quality, and Emotional Intelligence are the top-ranked predictors.

Figure 8 presents class-specific SHAP importance for each of the three outcome classes. For Cognitive Decline Risk, low Mental Wellbeing Score and low Social Interaction Quality are the primary positive contributors. For Positive Growth, high Emotional Intelligence and high Human Cognitive Agency are the dominant positive drivers. For Stagnant/Moderate, no single feature dominates, consistent with this class representing a heterogeneous intermediate state.

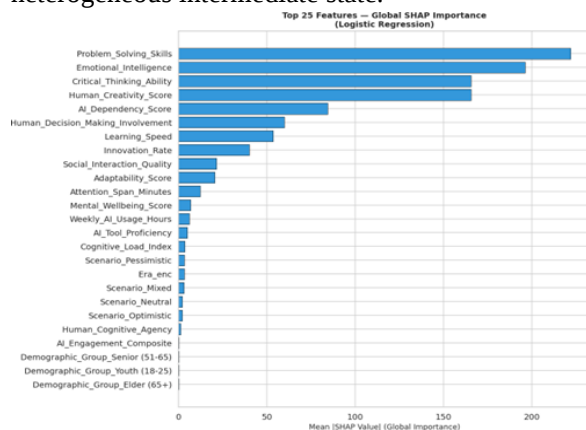


Figure 8. Class-specific SHAP feature importance for the three Outcome Label classes. Cognitive Decline Risk is driven by low wellbeing and social quality; Positive Growth by high emotional intelligence and human agency.

F. Categorical Analysis

Figure 9 presents outcome proportions across categorical variables (Region, Demographic Group, Education Level, Profession, Era, Scenario, AI Tool Category). Key patterns: the Pessimistic scenario has the highest pro-portion of Cognitive Decline Risk; Full AGI Integration tool users show a polarized distribution with both the highest Positive Growth and Cognitive Decline Risk pro-portions; and Youth (18–25) demographic cohorts show modestly higher Cognitive Decline Risk than Adult (36–50) cohorts.

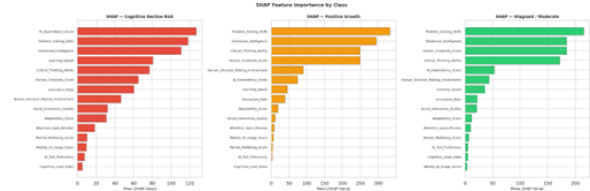


Figure 9. Outcome class proportions across categorical features. Pessimistic scenario and Full AGI Integration tool categories show the highest Cognitive Decline Risk proportions.

G. Discussion of Near-Perfect Performance

The uniformly high performance across all four models (accuracy > 98%) warrants careful interpretation. Several non-exclusive explanations are consistent with these results:

Linear separability of feature space:

The EDA reveals that three features (Mental Wellbeing Score, Social Interaction Quality, Emotional Intelligence) show strong, near-monotonic separation between the three classes. If the outcome variable was partially constructed from these features during dataset generation, a linear classifier would trivially achieve near-perfect separation.

Synthetic data generation artifacts:

The dataset spans hypothetical profiles across 2024–2075, suggesting it may be synthetically generated. Synthetically generated datasets often exhibit stronger feature-target associations than real-world data, leading to inflated model performance metrics.

Feature leakage risk:

While Positive Effects of AI and Negative Effects

of AI were excluded, the included psychocognitive features may be partially deterministic functions of the outcome label under the synthetic generation process.

We flag these as important interpretive caveats. The analytical pipeline and SHAP explainability framework remain valid and applicable to real-world data where performance metrics would likely be lower but the feature attribution insights would retain their directional validity.

VII. LIMITATIONS

L1 — Synthetic data provenance:

The dataset likely reflects a synthetic generation process. Causal and even associative claims cannot be generalized to real-world populations without replication on empirically collected data.

L2 — Cross-sectional structure: The data is cross-sectional despite spanning multiple years and eras. Longitudinal panel designs are required to establish temporal trajectories of AI-related cognitive change.

L3 — Excluded free-text features:

The Positive Effects of AI and Negative Effects of AI fields contain rich qualitative information that was excluded due to missingness and leakage risk. A proper NLP pipeline incorporating these fields may improve model interpretability and predictive diversity.

L4 — Potential data leakage:

Near-perfect classifier performance suggests possible information leakage between psychocognitive features and the outcome label under the synthetic generation process. Feature construction and target definition should be audited before any downstream policy application.

L5 — SHAP computational limitations:

SHAP values were computed on a 2,000-sample subsample of the test set for computational tractability. Full-test-set SHAP computation would provide more precise global importance estimates.

L6 — Hyperparameter optimization:

Models were trained with fixed hyperparameters.

Cross-validated grid search or Bayesian optimization may identify configurations with improved generalization to out-of-distribution data.

VIII. CONCLUSION

This study presents a systematic multiclass predictive analysis of AI-related mental growth trajectories across a rich 52,000-record dataset spanning global regions, demographic cohorts, and AI developmental eras. Four state-of-the-art classifiers were trained and evaluated under a rigorous stratified protocol with class-weighted training and 5-fold cross-validation.

Logistic Regression achieves the highest performance (test accuracy: 99.60%, macro F1: 0.9963, ROC-AUC: 0.9999), demonstrating that the class-discriminating signal in this dataset is substantially linear in structure. SHAP explainability analysis identifies Mental Wellbeing Score, Social Interaction Quality, and Emotional Intelligence as the dominant predictors of outcome class membership — a finding with direct implications for AI policy and digital mental health intervention design.

The central substantive conclusion is that the psychological and social dimensions of AI engagement, rather than raw usage intensity or technical proficiency, are the primary determinants of whether AI exposure results in positive cognitive growth or cognitive decline risk. This points toward wellbeing-centered and socially embedded AI deployment strategies, rather than the productivity-centric metrics that currently dominate AI policy discourse.

Future work should: (a) replicate this pipeline on empirically collected longitudinal data; (b) incorporate NLP-based processing of the free-text effects fields; (c) investigate causal mechanisms linking AI dependency and mental wellbeing using structural equation modeling; and (d) extend the analysis to domain-specific sub-populations (e.g., students, healthcare workers, creative professionals).

IX. ETHICAL CONSIDERATIONS

Predictive models of cognitive decline risk carry significant ethical stakes. Classification of

individuals into “Cognitive Decline Risk” categories based on AI usage patterns risks stigmatization, discriminatory insurance or employment decisions, and self-fulfilling prophecies if risk scores are disclosed without appropriate clinical context. Any deployment of such models in real-world settings must be accompanied by rigorous fairness auditing across demographic groups, transparent uncertainty quantification, and human oversight in consequential decision-making contexts.

DATA AND CODE AVAILABILITY

The dataset (AI Impact on Human Mental Growth.csv) and the complete analysis notebook (AI Impact Mental Growth Research.ipynb) are provided as supplementary materials. The notebook is Google Colabready and includes inline instructions for environment setup, data loading, preprocessing, model training, evaluation, and SHAP analysis.

AUTHOR CONTRIBUTIONS

Ranveer Singh, Sidhesh Mishra: Conceptualization, Methodology, Formal Analysis, Writing — Original Draft. Aryan Somanna, Riduvarshini M: Data Cu-ration, Visualization. Kandjoni Yendouname: Sta-tistical Analysis, Writing — Review. Shrey Jauhari, Venkatanathan O R V, Karthik H: Software, Writing — Review & Editing.

CONFLICT OF INTEREST

The authors declare no conflicts of interest.

REFERENCES

- [1] Brynjolfsson and A. McAfee, *The Second Machine Age: Work, Progress, and Prosperity in a Time of Brilliant Technologies*. New York, NY, USA: W. W. Norton & Company, 2014.
- [2] D. H. Autor, “Why Are There Still So Many Jobs? The History and Future of Workplace Automation,” *Journal of Economic Perspectives*, vol. 29, no. 3, pp. 3–30, 2015.
- [3] D. Acemoglu and P. Restrepo, “Artificial Intelligence, Automation, and Work,” NBER Working Paper No. 24196, National Bureau of Economic Research, 2018.
- [4] OECD, *OECD Employment Outlook 2023*:

Artificial Intelligence and the Labour Market. Paris, France: OECD Publishing, 2023.

- [5] J. M. Twenge, *iGen: Why Today’s Super-Connected Kids Are Growing Up Less Rebellious, More Tolerant, Less Happy and Completely Unprepared for Adulthood*. New York, NY, USA: Atria Books, 2017.
- [6] J. D. Mayer, P. Salovey, and D. R. Caruso, “Emotional Intelligence: Theory, Findings, and Implications,” *Psychological Inquiry*, vol. 15, no. 3, pp. 197–215, 2004.
- [7] A. B. R. Shatte, D. M. Hutchinson, and S. J. Teague, “Machine Learning in Mental Health: A Scoping Review of Methods and Applications,” *Psychological Medicine*, vol. 49, no. 9, pp. 1426–1448, 2019.
- [8] S. M. Lundberg and S.-I. Lee, “A Unified Approach to Interpreting Model Predictions,” in *Advances in Neural Information Processing Systems*, vol. 30, pp. 4765–4774, 2017.
- [9] L. Breiman, “Random Forests,” *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [10] T. Chen and C. Guestrin, “XGBoost: A Scalable Tree Boosting System,” in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 785–794, 2016.
- [11] G. Ke et al., “LightGBM: A Highly Efficient Gradient Boosting Decision Tree,” in *Advances in Neural Information Processing Systems*, vol. 30, pp. 3146–3154, 2017.