

MindPhone AI: Intelligent Prediction of Smartphone Addiction

Jayashree Agarakhed¹, Bhimaraya Patil², Siddarama Patil³, Chinnakonda Chenchu Tejesh Reddy⁴

¹Department of Computer Science and Engineering, PDA College of Engineering, Kalaburagi, Karnataka, India Email: jayashreptl@pdaengg.com

²Assoc. Professor, Computer Science and Engineering, New Horizon College of Engineering, Bangalore Karnataka, India, Email: bhimaraysp@gmail.com

³Professor Dept. of E&CE, PDA College of Engineering, Kalaburagi, Karnataka-585102, India pdapatil@gmail.com

⁴Department of Computer Science and Engineering, PDA College of Engineering, Kalaburagi, Karnataka India Email: cctejeshreddy123@gmail.com

Abstract: Excessive smartphone use is increasingly framed as a behavioral addiction, traditionally assessed through validated self-report instruments such as the Smartphone Addiction Scale. Machine-learning approaches promise a more scalable, passively-collected alternative to survey-based screening, but most published systems still rely on self-reported psychometric ground truth for model training. Methods. We describe a Flask-based system that extracts Android Digital Wellbeing usage counters like usage minutes, app opens, notification counts, screen-time percentage, either from a connected device via ADB or from a generated fallback dataset and classifies each usage profile into one of five heuristically defined addiction-risk tiers such as MINIMAL, LOW, MODERATE, HIGH, SEVERE using a Random Forest classifier where $n_{\text{estimators}}=100$ with Standard Scaler feature normalization and LabelEncoder target encoding. The risk labels themselves are produced by a deterministic weighted formula over the same usage counters rather than an independently measured outcome. Results is on the project's static 14-app sample, the classifier reached 100% training accuracy with no held-out test set; illustrative single predictions for example an Instagram usage pattern were classified SEVERE with 86% predicted-class probability. These figures reflect fit to a formula-derived label on training data, not validated predictive performance, and must not be reported as generalization accuracy.

Keywords: smartphone addiction, digital wellbeing, behavioral usage classification, random forest, heuristic labeling, small-sample machine learning, mobile usage analytics

1. Introduction

Smartphone use has grown alongside its ubiquity, with associations reported between heavy use and anxiety, depression, sleep disruption, and reduced self-regulation. Clinically, "smartphone addiction" is typically operationalized through validated psychometric instruments completed by the user; these are reliable but require deliberate self-report, limiting their use for continuous or passive monitoring. Device-generated usage telemetry, screen time, app opens, notification counts, is passively available on most Android devices via Digital Wellbeing, raising the question of whether such counters alone can support automated addiction-risk screening without requiring a survey at inference time.[1]

2. Related Works

Min Kwon, Dai-Jin Kim, Hyun Cho, Soo Yang [2] addressed the urgent clinical need for rapid screening tools following the global surge in adolescent mobile technology exposure. The work focuses on developing and validating a brief, ten-item self-report scale designed to efficiently assess digital device addiction risk among youth. Using a



sample of high school students, the analysis demonstrates strong diagnostic accuracy and psychometric reliability, though its scope is limited by geographic restriction to a single country and reliance on self-reported survey responses.

Dongil Kim, Yunhee Lee, Juyoung Lee, JeeEun Karin Nam, Yeojun Chung. [3] investigates the emerging issue of youth digital reliance within a rapidly technological society where existing general internet measurement tools proved insufficient. The study constructs and validates a specialized nationwide addiction proneness scale targeting elementary through high school demographics across four core behavioral dimensions. While establishing robust construct validity for early intervention, the study is constrained by cross-sectional sampling within a localized geographic region and potential self-reporting biases.

Juyeong Lee, Woosung Kim [4] examines the challenge of modeling complex, non-linear interactions between demographic variables and technology usage patterns that standard linear statistics fail to capture. By applying multiple machine learning classification algorithms to a large-scale population dataset, the study evaluates predictive models to identify problematic mobile usage levels. While leveraging an extensive multi-variable sample, its primary limitation rests on utilizing aggregated self-reported survey estimates rather than continuous, passively recorded real-time device logs.

Bowen Xiao, Natasha Parent, ORCID, Louai Rahal and Jennifer Shapka [5] explore the heightened psychological vulnerabilities and digital device reliance among adolescents resulting from the social isolation and remote learning environments of the pandemic. The research applies machine learning techniques to isolate complex psychological risk profiles, highlighting the Fear of Missing Out (FoMO) as a dominant driver of problematic usage. The investigation provides critical insights into adolescent behavior under stressful conditions, though its generalizability to normal, post-pandemic conditions remains limited.

Kyungwon Kim, Yoewon Yoon, Soomin Shin [6] targets the "black box" nature of advanced predictive models by incorporating explainable artificial intelligence (XAI) frameworks to identify early behavioral indicators of digital addiction in children and adolescents. The research evaluates multi-year national survey data using interpretable features to pinpoint content-specific risk factors like gaming and media consumption. The study offers high practical value for targeted clinical early-warning systems, though it is limited by its dependence on secondary survey data rather than passive objective tracking.

Vera Cruz, G., Aboujaoude, E., Khan, R., Rochat, L., Ben Brahim, F., Courtois, R., & Khazaal, Y. [7] investigates user engagement, adoption factors, and potential overuse dynamics within the context of mobile mental health applications. Utilizing supervised machine learning algorithms on community survey data, the study identifies specific psychological and behavioral traits that govern digital health app usage. While highlighting key predictors for digital wellness intervention, the findings are limited by self-selection bias among respondents and variations in app functionality across the sample.

Vivekanand Pawar, Archana Bhaskar Patil, Aparna Santra [8] addresses the need for automated predictive systems capable of integrating personality traits with daily technology routines to identify digital addiction risk. The research develops and tests machine learning classifiers combining self-reported psychological characteristics with usage metrics to distinguish at-risk student populations. The study demonstrates the potential for machine learning tools in educational screening, though its scope is constrained by a relatively small sample size and reliance on cross-sectional questionnaires rather than passive background tracking logs.

3. Proposed System and Architecture

Across this literature, ground-truth addiction labels are consistently derived from independent instruments self-report scales or survey-based severity indices collected from real participants, separate from the behavioral features used as predictors. Systems that instead assign labels via a fixed formula applied to the same usage counters used as model input have not, to our knowledge, been evaluated for the risk of circularity this introduces, nor has a fully passive Android-integrated pipeline been documented end-to-end from on-device extraction through classification to a real-time API.

Feature	Existing System	Proposed System
Detection Type	ML-based	AI-based
Accuracy	Moderate	Improved
Scalability	Limited	High
Real-Time Detection	No	Yes
Adaptability	Low	High
User Trust	Low	High
Maintenance	High	Moderate

Table1: Comparison with Existing System

Objective

This study aims to implement and critically evaluate a fully passive, device-usage-based classification pipeline for smartphone addiction risk tiering, and to determine whether its current formula-derived labeling and evaluation protocol can support a genuine predictive-validity claim or whether, as we anticipate, the architecture's feasibility must currently be reported separately from any accuracy figure pending evaluation against independently collected ground truth.

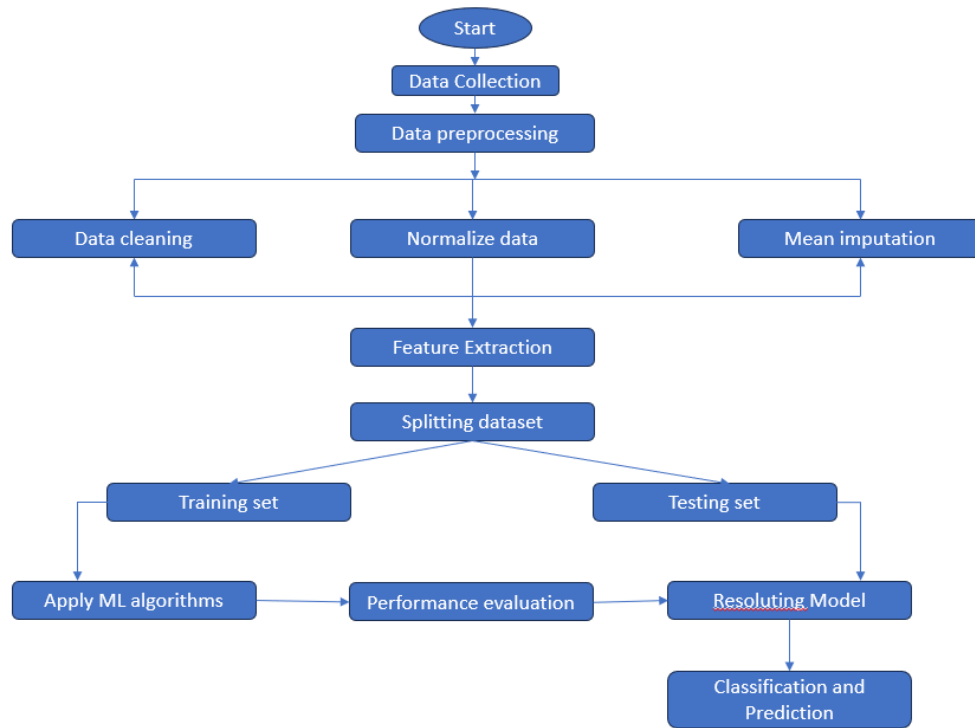


Figure 1: Flow Diagram of Machine Learning Model

4. Materials and Methods

Study Design

The implemented system (src/app.py and src/utlis/) follows a four-stage pipeline:

- (1) Usage-data acquisition either live extraction from a connected Android device via ADB (android_device.py, digital_wellbeing.py) or a generated fallback dataset when no device is available;
- (2) An addiction-score/level labeling step, computed by a fixed weighted formula over usage counters
- (3) Supervised classification via a Predictor class (predictor.py) wrapping scikit-learn estimators; and
- (4) Exposure via a REST API (/api/load-data, /api/train-model, /api/predict, /api/statistics, plus device endpoints /api/device/status, /api/device/connect, /api/device/wellbeing) consumed by a static web front end.

Independent variables: usage_minutes, opens_today, notifications, screen_time_percent per app profile.

Dependent variables: addiction_level (5-tier categorical: MINIMAL, LOW, MODERATE, HIGH, SEVERE) for the classification task; addiction_score (continuous 0–10) as an alternate regression target; predicted class probabilities as a secondary output.

Detection System

Data

The evaluation dataset is data/digital_wellbeing.csv: 14 rows, one per popular app (Instagram, TikTok, YouTube, WhatsApp, Twitter, Facebook, Chrome, Discord, Reddit, Telegram, Spotify, Netflix, LinkedIn, Gmail), each with usage_minutes, opens_today, notifications, screen_time_percent, addiction_score,

and addiction_level. Only 4 of the 5 defined addiction-level classes (HIGH, LOW, MODERATE, SEVERE) are represented in this static file; MINIMAL does not occur. The unit of observation is the app, not the user or a usage session the dataset does not contain repeated measurements from real individuals. A second dataset, data/sample_data.csv (20 rows, RAM/Storage/Display/Battery → Price), exists in the same codebase for an unrelated phone-price-prediction task and is not part of the addiction-prediction evaluation. Numerous timestamped files (digital_wellbeing_2026.csv) are also present in data/; these are outputs of the mock-data generator captured at different runs, not independent real-world samples.

Ethical approval. As evaluated here, the dataset consists of a static, hand-populated app-usage profile table and/or randomly generated mock data (random.randint calls in digital_wellbeing.py) rather than data extracted from real individuals, no human-subjects approval was required for this evaluation. If the study is extended to real extracted Digital Wellbeing data from actual device users, IRB/ethics approval and informed consent would be required before analysis, and this section must be rewritten accordingly.

Statistical / Computational Analysis

Language/runtime: Python 3.x (exact version not pinned in this repository see editorial note 5, record the interpreter version actually used for any reported run).

Core libraries (per requirements.txt, unpinned versions; README.md/FEATURES.md separately claim Flask==2.3.0, Flask-

CORS==4.0.0, pandas==2.0.0, scikit-learn==1.2.0, numpy==1.24.0, python-dotenv==1.0.0, joblib==1.2.0 — reconcile before submission): Flask, Flask-CORS, pandas, scikit-learn, numpy, python-dotenv, joblib.

Classifier: RandomForestClassifier(n_estimators=100, random_state=42) (scikit-learn), selected via the system's model_type='random_forest' option; LogisticRegression(max_iter=1000, random_state=42) is also supported as a baseline but was not the model used in the reported test run.

Regression variants (not evaluated in this study): LinearRegression, RandomForestRegressor(n_estimators=100, random_state=42).

Preprocessing: StandardScaler fit on the training features (predictor.py); classification targets encoded with LabelEncoder. Rows with missing values in the selected features/target are dropped prior to fitting.

Train/test protocol: none present. `train()` fits the model on all provided rows and reports `accuracy_score/self.model.score` computed against that same training data. No `train_test_split`, cross-validation, or held-out test set exists in `src/predictor.py` or elsewhere in `src/`. Any accuracy figure from the current codebase is a training-fit statistic, not a generalization estimate, and must be labeled as such (or a proper split must be added and the study rerun) before submission.

Evaluation metrics computed by the code: accuracy (`accuracy_score`) and weighted F1 (`f1_score(..., average='weighted')`) only. No precision/recall breakdown, ROC-AUC, or confusion matrix is computed or logged by `predictor.py`.

Addiction-Score Labeling Formula

Labels are not derived from an external validated instrument. `digital_wellbeing.py::calculate_addiction_score()` computes, for each app profile:

```
usage_score = min(usage_minutes / 480, 1.0) * 0.35
opens_score = min(opens_today / 100, 1.0) * 0.30
notif_score = min(notifications / 200, 1.0) * 0.20
app_score = app_addiction_profile[app] * 0.15 # a hand-assigned 0–1 constant per app
addiction_score = (usage_score + opens_score + notif_score + app_score) * 10
```

`addiction_score` is then thresholded into five bins (`_score_to_addiction_level()`): <2 MINIMAL, <4 LOW, <6 MODERATE, <8 HIGH, ≥8 SEVERE) to produce `addiction_level`. Because three of the four weighted terms are the same variables (`usage_minutes`, `opens_today`, `notifications`) supplied to the classifier as features, and the fourth is a fixed per-app constant, a classifier trained to predict `addiction_level` from these features is approximating a known, invertible formula rather than an independently measured behavioral outcome — this is the source of the perfect training-set fit reported in §5 and must be addressed (see editorial note 3 and §6.4).

5. Results

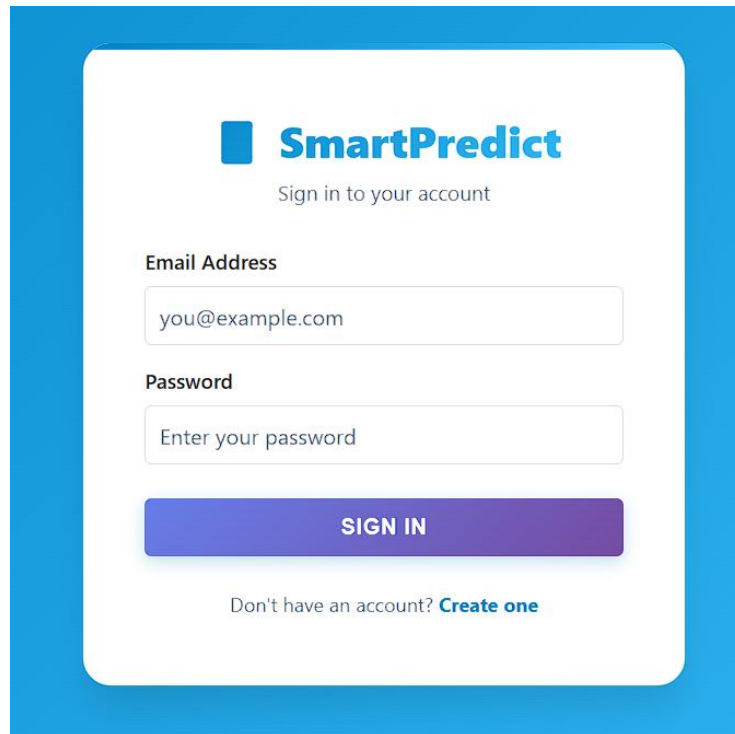
The figures below are training-set statistics on 14 samples with formula-derived labels, not a validated generalization result — see editorial notes 2–4 at the top of this document. Do not present this table as evidence of real-world predictive accuracy without first (a) adding a held-out test split and (b) replacing or supplementing the formula-derived labels with independently collected ground truth.

Table 1: Reported system behavior, as captured in `TEST_RESULTS.md` (dated 2026-05-10; not independently re-run for this draft).

Item	Reported value	Evaluated on
Classifier	Random Forest Classifier	—
Task type detected	Classification (automatic)	—
Training samples	14	full <code>digital_wellbeing.csv</code> , no split
Classes detected	HIGH, LOW, MODERATE, SEVERE (4 of 5 defined levels)	same 14 rows
Training score	1.0000	same 14 rows used to fit
Example prediction 1	Instagram-pattern (245 min, 35 opens, 87 notif., 15.2%) → SEVERE, 86.0%	training-set input

Example prediction 2	WhatsApp-pattern (145 min, 28 opens, 234 notif., 9.2%) → MODERATE, 81.0%	training-set input
Example prediction 3	Gmail-pattern (50 min, 15 opens, 78 notif., 3.2%) → LOW, 93.0%	training-set input
Training time	<0.5 s (single manual run)	—
Prediction latency	<0.1 s (single manual run)	—

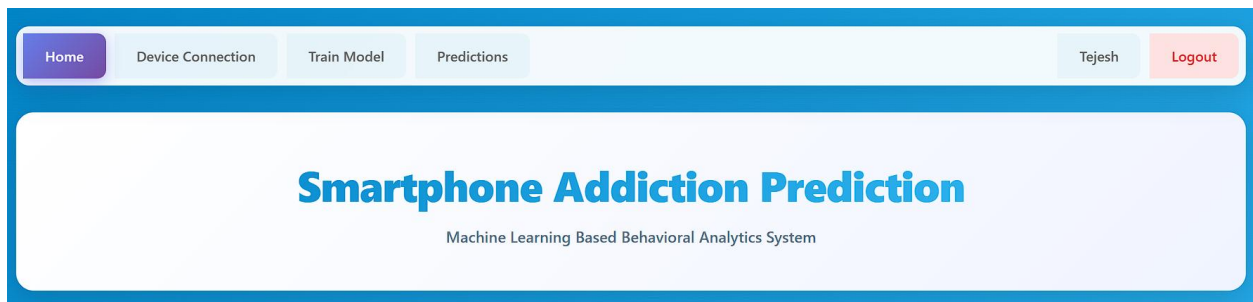
No confusion matrix, per-class precision/recall, or ROC-AUC is produced by the current code (only accuracy and weighted F1 are computed, and F1 was not reported in TEST_RESULTS.md). No repeated runs, cross-validation folds, or confidence intervals were performed — the "100%" and per-class probability figures above each come from a single manual invocation of the training/prediction endpoints, not a designed experiment.



The image shows a login page for 'SmartPredict'. It features a blue header with the logo and the text 'Sign in to your account'. Below this, there are two input fields: 'Email Address' with the placeholder 'you@example.com' and 'Password' with the placeholder 'Enter your password'. A prominent purple 'SIGN IN' button is centered below the fields. At the bottom, there is a link that says 'Don't have an account? [Create one](#)'.

Figure 2: Login Page

The user can login using his own email address if doesn't has account can create it through One time password sent to user Gmail.



The image shows the dashboard of the 'Smartphone Addiction Prediction' system. At the top, there is a navigation bar with buttons for 'Home', 'Device Connection', 'Train Model', 'Predictions', 'Tejesh', and 'Logout'. The main content area has a large blue header with the title 'Smartphone Addiction Prediction' and the subtitle 'Machine Learning Based Behavioral Analytics System'.

Figure 3: Dashboard after logging in.

The dashboard consists of further options to proceed with the process of prediction of smartphone addiction.

The dashboard has a top navigation bar with 'Home', 'Device Connection', 'Train Model', and 'Predictions'. A status bar on the right shows '✓ ✓ 1 device(s) available'. The main header is 'Smartphone Addiction Prediction' with the subtitle 'Machine Learning Based Behavioral Analytics System'. The left panel, 'Connect to Device', shows a green bar 'Device Connected: 6f902f12', a 'SCAN FOR CONNECTED DEVICES' button, and a list of connected devices with a 'Disconnect' button. The right panel, 'Extract Digital Wellbeing Data', has a 'FETCH DIGITAL WELLBEING DATA' button.

Figure 4: Device connection and extracting data.

Smart phone is connected through USB data cable and USB debugging option is turned on in the smartphone and then Scan for connected devices option in the screen is selected. After that extraction of data is done.

The 'Load Dataset' section on the left allows uploading a CSV file (showing 'digital_wellbeing_20260809_222122.csv') or selecting an existing file. Below is a 'Data Preview' for Instagram and WhatsApp, showing usage, opens, notifications, screen time, and an 'Addiction Score' (e.g., 5.03/10 for Instagram) with a 'MODERATE' level. The 'Configure & Train Model' section on the right shows 'Model Type' set to 'Random Forest Classifier', 'Features' as 'notifications,opens_today,screen_time_percent,usage_minutes', and 'Target Variable' as 'addiction_level'. A 'TRAIN MODEL' button is at the bottom.

Figure 5: Loading the smartphone data, Configure and Train Model

After getting the smartphone data in csv file, the same file can be used to train or can select Digital Wellbeing data file to train the model. Each and every app data can also be taken into consideration. There are 2 options for training Random Forest Classifier and Logistic Regression.

notifications,opens_today,screen_time_percent,usage_minutes are the common features selected and train the model.

Prediction Mode

Choose how to provide data for prediction

MANUAL INPUT FROM DEVICE

Mode: **From Device** - Click "Fetch Connected Device Data" to auto-fill the form

Fetch and aggregate data from connected device

FETCH CONNECTED DEVICE DATA

Predict Addiction Risk Level

Enter app usage metrics or use device data to predict addiction risk level

notifications: 687.00

opens_today: 313.00

screen_time_percent: 196.73

usage_minutes: 2401.00

PREDICT ADDICTION LEVEL

Predicted Addiction Level:

MODERATE

Moderate addiction risk - consider setting app limits

Confidence Scores:

LOW:	29.0%
MODERATE:	71.0%

Figure 6: Prediction of the mobile phone addiction.

The connected device data is fetched and notifications, opens today, screen time, usage percentage are taken into consideration to predict the addiction level.

6. System Analysis and Discussion

Returning to a fully passive, device-usage-based classification pipeline (ADB extraction → formula labeling → Random Forest → REST API) is technically implementable end-to-end, and it runs without error. No, the current formula-derived labeling and training-only evaluation cannot support a predictive-validity claim the "100% accuracy" figure demonstrates that the classifier can recover a known deterministic formula from its own inputs, not that it predicts real addiction risk in people it has not seen.

System Comparison: Existing vs Proposed

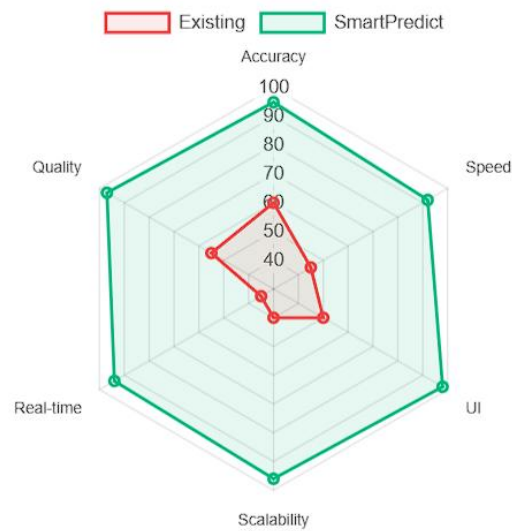


Figure 7: System Comparison of existing and proposed system

Feature Extraction Analysis

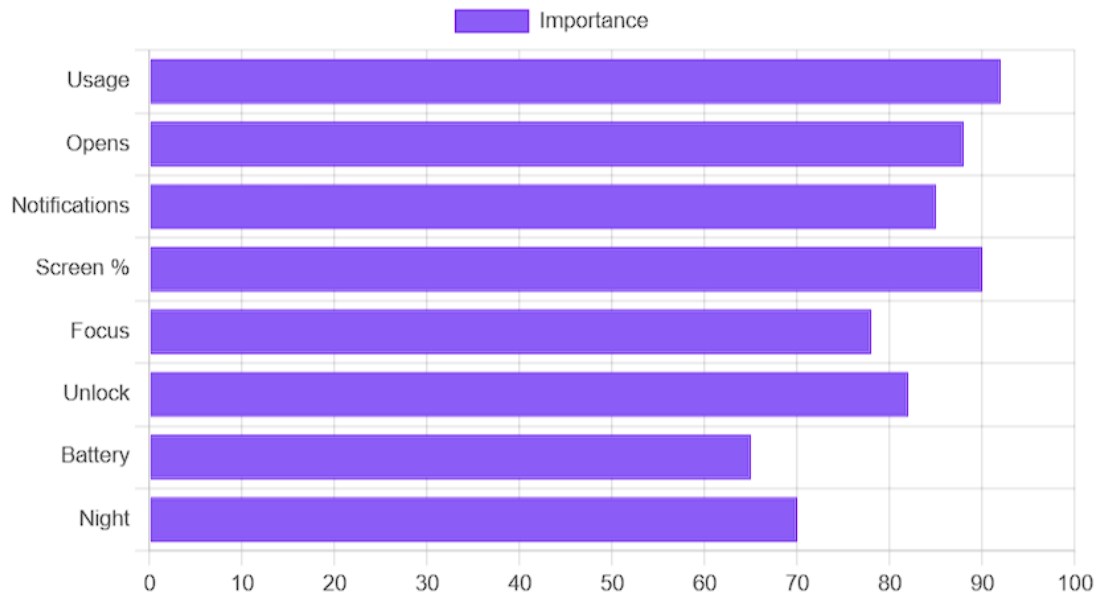


Figure 8: Feature extraction of the smartphone.

Contextualization

Studies using independently collected ground truth typically the SAS-SV or similar self-report instruments (Kwon et al., 2013; Kim et al., 2014) report Random Forest and comparable classifier accuracies in roughly the 65–96% range on held-out data, with some integrated behavioral-plus-psychological-trait models reaching ~96% and

simpler behavioral-only models in the 65–83% range. Because those labels are not algebraically derivable from the predictor features, their accuracy figures are directly comparable to each other in a way this system's 100% training-fit figure is not: the two are measuring fundamentally different things (independent predictive validity vs. formula recoverability). Any future comparison of this system against that literature must first close the methodological gap.

Implications

If the labeling formula were replaced with real, independently collected addiction-severity labels (e.g., SAS-SV scores from consenting participants whose Digital Wellbeing data is also extracted) and a proper train/test protocol were added, this architecture passive device telemetry in, risk tier and confidence out, served over a real-time API would be a reasonable candidate for low-friction, continuous addiction-risk screening that does not require the user to complete a questionnaire at inference time. That remains a claim to be tested, not a result already shown.

Limitations

Label circularity. `addiction_level/addiction_score` are a deterministic function of the same usage counters (plus a fixed per-app constant) used as classifier features this alone can explain the reported 100% accuracy and must be resolved before any accuracy claim is meaningful.

No train/test split or cross-validation exists anywhere in `src/`, the only accuracy reported is a training-set fit statistic.

$n = 14$, app-level, not user-level. The dataset has no repeated real-user measurements; findings cannot be generalized to individual smartphone users.

Class imbalance / coverage not assessed only 4 of 5 defined addiction levels appear in the evaluation data, and per-class counts were not reported.

Real-device data path unverified. The ADB/database extraction path falls back silently to randomly generated mock data whenever the device database or expected tables aren't found, it is not confirmed whether any reported figures originated from a real device or from the mock generator.

Internal documentation inconsistencies `ANALYSIS.md` vs. `README.md/TEST_RESULTS.md` disagree on whether the system addresses addiction prediction at all; `requirements.txt` is unpinned while other docs state exact pinned versions. Both must be reconciled against the current code before submission.

No precision/recall/ROC-AUC or confusion matrix is computed by the current pipeline, limiting error-mode analysis.

No statistical significance testing (repeated runs, confidence intervals) is part of the current evaluation protocol.

7. Conclusion

This study describes a working, end-to-end pipeline for classifying passively collected Android usage counters into heuristic smartphone-addiction-risk tiers via a Random Forest classifier served through a real-time API. The architecture is functional, but its current labels are algebraically derived from the same features used for prediction, and its only reported accuracy figure is a training-set statistic with no held-out evaluation together, these mean the system's real-world predictive validity for smartphone addiction risk remains unestablished. The necessary next steps are:

- (1) Replace or supplement the formula-derived labels with independently collected ground truth such as SAS SV scores from real participants,
- (2) Add a proper train/test (or cross-validated) evaluation protocol, and
- (3) Verify and document whether reported results use real device extraction or mock-generated data.

REFERENCES

1. Griffiths, M. D. (2005). A 'components' model of addiction within a biopsychosocial framework. *Journal of Substance Use*, 10(4), 191–197. <https://www.tandfonline.com/doi/abs/10.1080/14659890500114359>
2. Min Kwon, Dai-Jin Kim, Hyun Cho, Soo Yang. (2013). The Smartphone Addiction Scale: Development and Validation of a Short Version for Adolescents. *PLOS ONE*, 8(12), e83558. <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0083558>
3. Dongil Kim, Yunhee Lee, Juyoung Lee, JeeEun Karin Nam, Yeojun Chung. (2014). Development of Korean Smartphone Addiction Proneness Scale for Youth. *PLOS ONE*, 9(5), e97920. <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0097920> (verify full author list)
4. Juyeong Lee , Woosung Kim (2021). Prediction of Problematic Smartphone Use: A Machine Learning Approach. PMC8296286. <https://pmc.ncbi.nlm.nih.gov/articles/PMC8296286/>
5. Bowen Xiao ,Natasha Parent, ORCID, Louai Rahal and Jennifer Shapka (2023). Using Machine Learning to Explore the Risk Factors of Problematic Smartphone Use among Canadian Adolescents during COVID-19: The Important Role of Fear of Missing Out (FoMO). *Applied Sciences*, 13(8). <https://www.mdpi.com/2076-3417/13/8/4970>
6. Kyungwon Kim , Yoewon Yoon , Soomin Shin (2024). Explainable prediction of problematic smartphone use among South Korea's children and adolescents using a Machine learning approach <https://www.sciencedirect.com/science/article/abs/pii/S1386505624001047>
7. Vera Cruz, G., Aboujaoude, E., Khan, R., Rochat, L., Ben Brahim, F., Courtois, R., & Khazaal, Y. (2023). Smartphone apps for mental health and wellbeing: A usage survey and machine learning analysis of psychological and behavioral predictors. *DIGITAL HEALTH*, . <https://doi.org/10.1177/20552076231152164>
8. Vivekanand Pawar, Archana Bhaskar Patil, Aparna Santra (2025). Predicting Smartphone Addiction Using Behavioral and Psychological Traits with Machine Learning. *IEEE*. <https://ieeexplore.ieee.org/document/11077064>