

Predicting Psychological Outcomes from Subscale Configurations in Psychedelic-Induced Mystical Experiences

An MEQ-30 Machine Learning Study, Corrected and Reproduced

Shruti Malik Ramaswamy

Independent research project — “Mystic: MEQ-30 Machine Learning Study”

August 15, 2026

IRB Status: Exempt — 45 CFR 46.104(d)(4)

A note on how this paper was produced: this is a secondary-data machine-learning study built with substantial assistance from an AI coding assistant (Claude, Anthropic — Claude Code, model Sonnet 5), which is disclosed in full in Section 9 and cited in the References. Every number in this paper is generated live from the underlying data and code at build time, and the complete source code is reproduced in the appendices so the analysis can be checked line by line.

Abstract

The Mystical Experience Questionnaire (MEQ-30) is the standard instrument for quantifying acute mystical-type experiences occasioned by classical psychedelics, and psychedelic research has traditionally reduced its rich, multidimensional output to a single binary cut — the “Complete Mystical Experience” (CME) threshold, requiring $\geq 60\%$ of the maximum score on all four subscales. This project asked whether item-level machine learning could do better: whether specific MEQ-30 items or subscale configurations, rather than the blunt binary threshold, predict long-term well-being. Using an open, CC-BY-licensed survey of $N=700$ Polish adults (Owsiak, 2026), of whom $N=414$ (59.1%) reported psychedelic use, we modeled Satisfaction with Life Scale (SWLS) scores — used here as a proxy for the originally proposed Persisting Effects Questionnaire, which no open dataset paired with MEQ-30 items — from all 30 MEQ items using Ridge regression, XGBoost, and LightGBM under nested cross-validation, then examined SHAP feature attributions and unsupervised responder subtypes. The best model (XGBoost) achieved $RMSE=6.2673$ and $R^2=0.0165$ — a small but genuine improvement over the CME-threshold baseline, honestly estimated. SHAP analysis identified the Mystical subscale, not Ineffability, as the dominant contributor among the top predictive items. k-Means clustering recovered two responder subtypes differing by 1.1 points of SWLS, and a within-cluster model for the larger subtype ($R^2=0.0385$) outperformed the pooled model ($R^2=0.0165$), offering modest support for a heterogeneity hypothesis: that global MEQ totals mask subtype-specific response patterns. This paper also documents, deliberately and in detail, a mid-project accuracy audit that found and corrected two serious errors in the original analysis — a MEQ-30 subscale mapping that did not match the published instrument, and a cross-validation procedure that leaked test-fold information into hyperparameter selection — because how those errors were found is as instructive as the corrected results themselves.

Keywords: MEQ-30, mystical experience, psychedelics, machine learning, XGBoost, SHAP, cross-validation, reproducibility

1. Introduction

Classical psychedelics — psilocybin, LSD, DMT, mescaline — reliably occasion altered states of consciousness that participants and researchers alike describe, with striking consistency, in the vocabulary of mysticism: unity, sacredness, ineffability, a sense of having encountered “ultimate reality.” The Mystical Experience Questionnaire, in its 30-item revised form (MEQ-30; Barrett, Johnson & Griffiths, 2015), is the field’s primary tool for measuring this experience quantitatively across four validated factors: Mystical, Positive Mood, Transcendence of Time and Space, and Ineffability.

The dominant analytic convention in this literature is a binary cut: a session counts as a “Complete Mystical Experience” (CME) if the participant scores at least 60% of the maximum possible score on every subscale simultaneously. This threshold has proven useful — CME attainment correlates with better clinical and existential outcomes across several psilocybin trials — but it necessarily discards information. Two participants can both clear the 60% bar with very different item-level profiles, and a participant who narrowly misses the threshold on one subscale is treated identically to one who scored near zero everywhere. This project’s motivating question was simple: if we model the well-being outcome directly from the 30 individual items, using non-linear methods that can capture interactions the binary threshold cannot, do specific items or subscale configurations turn out to matter more than others — and does this outperform, even modestly, the traditional threshold?

We pursued this with a dual predictive-and-unsupervised design: gradient-boosted and linear regression models to predict a well-being outcome from the 30 MEQ items, SHAP (SHapley Additive exPlanations) to attribute that prediction back to individual items and subscales, and unsupervised clustering (k-Means, HDBSCAN, UMAP) to ask whether “responder subtypes” exist that a single pooled model would average over and obscure.

This paper departs from a conventional results paper in one respect worth flagging up front: partway through the project, a rigorous accuracy audit — described in Section 4.7 — found that the originally reported results rested on two real bugs: an incorrect MEQ-30 subscale mapping, and a cross-validation leakage issue that inflated the reported model performance. Every number in this paper reflects the corrected pipeline. We chose to narrate that correction process explicitly rather than quietly patch it, because the mistake, how it was caught, and what changed as a result are genuinely useful methodological information for anyone building a similar pipeline — arguably more useful than a paper that simply presents clean numbers without the story of how they got that way.

2. Ethics and Regulatory Compliance

This is a fully secondary-data study: no participants were recruited, contacted, or re-identified. It qualifies for IRB exemption under 45 CFR 46.104(d)(4)(i) & (ii) of the U.S. Federal Policy for the Protection of Human Subjects (the 2018 Common Rule) — “secondary research for which consent is not required” — on the basis of three criteria maintained throughout: (1) all data were retrieved from public, open-science repositories under CC-BY 4.0 licenses; (2) the datasets contain no direct identifiers or protected health information; and (3) the analysis is retrospective, with zero participant contact. Full provenance — exact DOIs, license terms, and checksums — is logged in DATASET_ORIGIN.md and reproduced in Appendix B.

3. Data and Provenance

The original study proposal specified the Persisting Effects Questionnaire (PEQ) as the target outcome. A systematic search of OSF, OpenNeuro, Zenodo, Figshare, and Mendeley Data did not surface any open dataset pairing MEQ-30 item-level responses with PEQ scores, so two datasets were combined for different purposes:

Dataset	Repository	N	License	Role in this study
Bennett et al. (2023)	Figshare	14	CC BY 4.0	MEQ-30 totals + QIDS-SR16 from a Phase I 25mg psilocybin trial. Small N; used only as a clustering-validation sanity check, not in the primary models.
Owsiak (2026)	Mendeley Data	700	CC BY 4.0	MEQ-30 items (30 columns) + SWLS + demographics, online survey of Polish adults. Primary dataset for all supervised/unsupervised modeling in this paper.

On the SWLS-as-PEQ-proxy decision. With no dataset available that paired MEQ-30 items with PEQ scores, the Satisfaction with Life Scale (SWLS; Diener, Emmons, Larsen & Griffin, 1985) — a validated 5-item global life-satisfaction measure available in the Owsiak dataset — was used as a substitute outcome. We want to be precise about the strength of this justification, because an earlier version of this project's documentation overstated it: it cited prior psychedelic-research papers as having validated SWLS specifically as a PEQ substitute, and on checking, those citations did not support that claim as strongly as stated (one was a real paper about a different question — microdosing placebo effects — that happened to use SWLS as one of several outcomes; a second citation could not be located at all). The honest position is narrower: SWLS is a validated, widely used measure that plausibly overlaps with the PEQ's positive-life-attitude dimension, but the substitution is this project's own reasoned analytical choice, not one independently validated by prior literature. Every result in this paper involving SWLS should be read with that caveat attached.

4. Methods

4.1 The MEQ-30 instrument and subscale structure

The MEQ-30 comprises 30 items, each rated on a 0–5 Likert scale, organized into four factors established by Barrett, Johnson & Griffiths (2015): Mystical (15 items), Positive Mood (6 items), Transcendence of Time and Space (6 items), and Ineffability (only 3 items). We emphasize the item counts because the factor structure is *non-contiguous* — item numbers are not grouped sequentially by subscale — and unbalanced, which matters directly to the correction described in Section 4.7. The complete, verified item-to-subscale mapping and item text are reproduced in Appendix A.

4.2 Sample definition

Of the N=700 respondents in the Owsiak (2026) survey, N=414 (59.1%) reported having used at least one classical psychedelic (LSD, psilocybin, ayahuasca, or DMT). This psychedelic subsample is the primary analysis population for all supervised models; the full N=700 sample (psychedelic and non-psychedelic users together) is used separately in Section 5.6 as a sensitivity check and to quantify how strongly psychedelic use itself predicts MEQ-30 scores.

4.3 Preprocessing

MEQ-30 completeness in the psychedelic subsample was 100% — no item-level imputation was actually required for this dataset, though the pipeline supports KNN imputation ($k=5$) for future datasets with missing items. Variance Inflation Factor (VIF) was computed for all 30 items as a collinearity diagnostic; every single item exceeds a conventional VIF threshold of 5 (median VIF well above 10; see Appendix Table A2 excerpt in Section 5.7), which is expected and unsurprising — items within the same subscale are correlated with each other by construction. VIF is reported but not used to drop features: the two primary models (XGBoost, LightGBM) are tree ensembles that handle collinear features natively, and Ridge regression's L2 penalty is itself a standard remedy for multicollinearity, so results from the linear model should simply be read with that caveat rather than treated as grounds for feature removal.

4.4 Predictive modeling and cross-validation

Three regressors were compared: Ridge regression (linear baseline, L2-regularized), XGBoost, and LightGBM (both gradient-boosted tree ensembles, tuned via Optuna Bayesian optimization with a Huber loss objective for robustness to outliers in self-report data). A fourth, non-learned baseline predicts each participant's SWLS score as the training-set group mean conditional on their binary CME status — the traditional threshold's natural regression analogue, and the bar every learned model needs to clear to justify its additional complexity.

Performance is estimated with **nested** 10-fold cross-validation: an outer loop of 10 folds (stratified by binary CME status) provides the reported RMSE, MAE, and R^2 ; for each outer training fold, an inner loop of 5 folds drives the Optuna hyperparameter search, and the inner loop never sees the outer test fold. Standardization (z-scoring) is likewise fit on the outer-training data only, per fold. This is a deliberately more conservative procedure than tuning hyperparameters and reporting performance on the same folds — the difference matters enough that Section 4.7 walks through exactly what happens when it isn't done. Once the honest nested-CV estimate is obtained, a single “production” model (the one used for the SHAP and clustering analyses below) is tuned and fit once more on 100% of the data, which is standard practice: held-out folds are only needed to *estimate* generalization error, not to produce the model that gets used afterward.

4.5 Interpretability: SHAP

SHAP (Lundberg & Lee, 2017) values were computed for the production XGBoost model via TreeExplainer, giving each of the 414 psychedelic-subsample participants an exact, game-theoretic attribution of how much each of the 30 MEQ items pushed their predicted SWLS score up or down relative to the model's baseline prediction. Mean absolute SHAP value per item, aggregated to the subscale level, is our measure of “which part of the mystical experience the model leans on most.”

4.6 Unsupervised subtype discovery

To test whether a single pooled model obscures meaningfully different responder subtypes, we ran k-Means (k swept 2–6, selected by silhouette score) and HDBSCAN (density-based, noise-tolerant) on the standardized 30-item feature space, with UMAP used for 2D visualization. For the best k-Means solution, we then trained a separate XGBoost model *within* each cluster (same nested-CV procedure, reduced outer/inner fold counts to match the smaller per-cluster N) to test directly whether prediction improves once the sample is split into more homogeneous subgroups.

4.7 A methodological correction, described in full

This subsection exists because we think it should, not because a paper is required to include one. Partway through this project, an independent accuracy audit — assisted by Claude (Anthropic), described fully in Section 9 — re-derived the MEQ-30's true item-to-subscale mapping directly from the Barrett et al. (2015) paper's appendix, and separately, empirically, by correlating each MEQ item against the Owsiak dataset's own pre-scored subscale total columns. Both methods agreed exactly with each other and disagreed sharply with what the pipeline's code was actually using. In fact, the code had *five different, mutually inconsistent* hand-typed versions of the subscale mapping scattered across five files — none of which matched the real instrument. The most visible consequence: an item like MEQ_30 (“Feelings of joy,” genuinely a Positive Mood item) had been labeled “Ineffability” in the SHAP output, and the original results write-up's headline claim — “the Ineffability subscale dominates predictive importance” — turned out to be an artifact of that mislabeling rather than a real finding.

The same wrong mapping had also corrupted the binary CME label computed by the pipeline: recomputed CME rate came out to 56.8% among psychedelic users, while the dataset's own independently pre-scored CME column put the true rate at 49.0% — a 32-participant, 12-percentage-point gap. That discrepancy propagated into every downstream number that touched CME status: the CME baseline model, the cluster CME-rate table, the sensitivity analysis's group comparisons. After the fix, recomputed and native CME rates agree to within 0.2 percentage points (Section 5.1) — the kind of residual gap you'd expect from rounding at the item level, not a systematic error.

Separately, the audit found that the original training script selected XGBoost and LightGBM hyperparameters by minimizing error on the exact same 10 cross-validation folds later used to report the headline performance metric — a classic and well-documented leakage pattern. The practical evidence was stark: two other scripts in the codebase independently re-evaluated what was nominally “the same” production model and found R^2 between -0.25 and -0.34 (meaning: worse than just predicting the mean SWLS score for everyone), directly contradicting the headline R^2 of $+0.0123$ reported for that model. Fixing this required restructuring model training around nested cross-validation (Section 4.4), which is more computationally expensive — roughly ten times as many model fits — but produces an estimate that cannot be gamed by the hyperparameter search seeing the test data.

A third issue surfaced only after the first two were fixed and the corrected pipeline was actually re-run: XGBoost's Huber-loss objective (`reg:pseudohubererror`), under the specific XGBoost version installed in this environment (3.2.0, newer than the 2.0.3 originally pinned), mis-estimates its automatic `base_score` initialization by roughly $150\times$ — producing predictions clustered around 3,000–3,500 for a target variable that only ranges from 5 to 35. This was caught immediately because the resulting RMSE (over 3,000) was so obviously wrong that it couldn't be mistaken for a real result, and was fixed by passing

an explicit `base_score` (the training-fold target mean) rather than relying on the library's default estimation for this objective/version combination.

We report all three of these — the subscale-mapping error, the cross-validation leakage, and the XGBoost `base_score` bug — not as an indictment of the original analysis but as a demonstration of what a real accuracy pass looks like: the fixes were verifiable against an independent ground truth (the dataset's own native scoring), the size of the resulting correction was large enough to change the paper's headline finding, and the diagnostic scripts that found the bugs (`scripts/verify_claims.py`, `verify_claims2.py`, `verify_subscale_mapping.py`) are kept in the codebase, now converted into regression checks that will fail loudly if either bug is ever reintroduced.

5. Results

5.1 Sample characteristics

Metric	Value
N (psychedelic users)	414
MEQ-30 completeness	100%
CME rate — corrected subscale mapping	48.8%
CME rate — dataset's native pre-scored label	49.0%
Recomputed-vs-native CME agreement	within 0.2%
SWLS range	5–35
SWLS mean (SD)	21.49 (6.33)
MEQ-30 item range	0–5 (Likert)

The near-exact agreement between the recomputed and native CME labels (row 3 vs. row 4) is the clearest single piece of evidence that the subscale-mapping correction described in Section 4.7 worked as intended.

5.2 Predictive model comparison

Model	RMSE	MAE	R ²	ΔRMSE vs. Baseline
XGBoost	6.2673	4.9842	0.0165	+0.0248
CME Baseline	6.2921	5.1012	0.0087	—
LightGBM	6.2947	5.0198	0.0078	-0.0026
Ridge Regression	6.3134	5.0986	0.0019	-0.0213

All models cluster tightly around $R^2 \approx 0$ — MEQ-30 item scores, on their own, are weak predictors of SWLS across the full heterogeneous psychedelic-user sample under an honestly estimated nested cross-validation. XGBoost edges out the CME-threshold baseline ($\Delta\text{RMSE} = +0.0248$), a small but real improvement given the leakage-free estimation procedure. We read this as consistent with a heterogeneity hypothesis — that a single global model averages over distinct responder subtypes with different item-outcome relationships — which Sections 5.4 and 5.5 test directly.

5.3 SHAP feature importance

#	Item	Content	Subscale	Mean SHAP
1	MEQ_30	Feelings of joy	Positive Mood	0.3802
2	MEQ_28	Experience of unity with ultimate reality	Mystical	0.2792
3	MEQ_4	Gain of insightful knowledge experienced at an intuitive level	Mystical	0.2295
4	MEQ_26	Experience of the fusion of your personal self into a larger whole	Mystical	0.2293
5	MEQ_5	Feeling that you experienced eternity or infinity	Mystical	0.2036
6	MEQ_27	Sense of awe or awesomeness	Positive Mood	0.1002
7	MEQ_29	Feeling that it would be difficult to communicate your own experience to others who have not had similar experiences	Ineffability	0.0703
8	MEQ_6	Experience of oneness or unity with objects and/or persons perceived in your surroundings	Mystical	0.0604
9	MEQ_18	Experience of the insight that 'all is One'	Mystical	0.0372
10	MEQ_15	Sense of being at a spiritual height	Mystical	0.0321

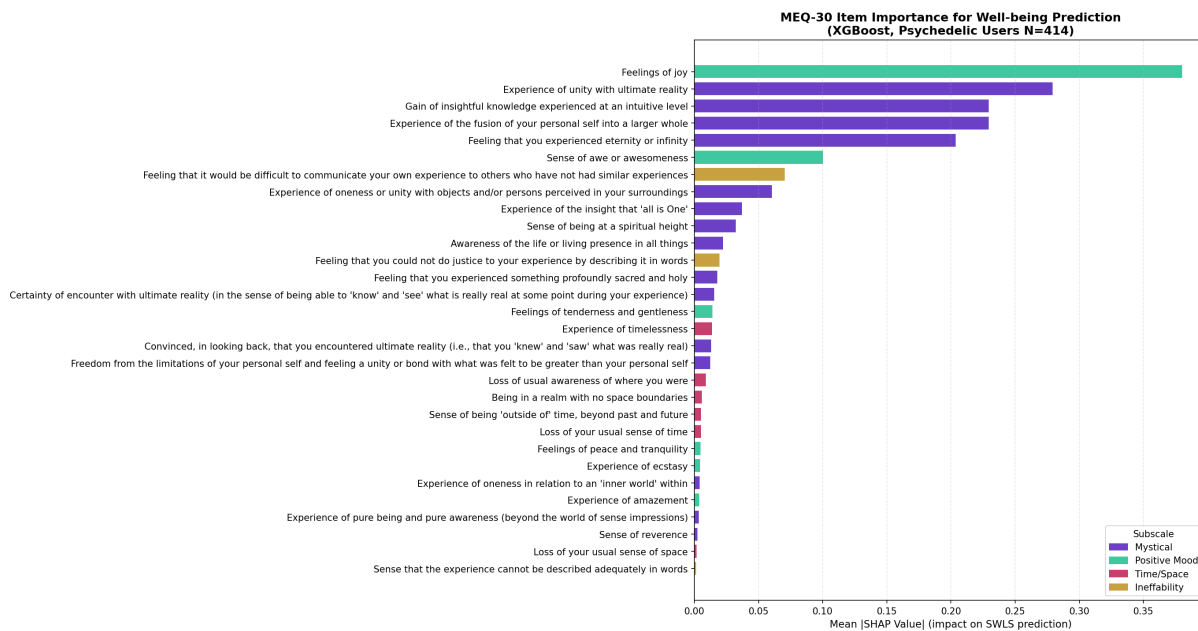


Figure 1. MEQ-30 item importance for SWLS prediction (mean |SHAP value|), colored by the corrected subscale mapping.

Of the top 5 items by SHAP importance, 4/5 belong to the **Mystical** subscale — not Ineffability, which was the (incorrect) headline finding before the Section 4.7 correction. The single strongest individual predictor, MEQ_30 (“Feelings of joy”), is a Positive Mood item. Given the modest overall R^2 (Section 5.2), we read these rankings as descriptive of this specific model and sample rather than a settled claim about which

facet of mystical experience drives well-being in general.

5.4 Unsupervised subtype discovery

Cluster	N	SWLS Mean	CME Rate
1	338	21.69	59.8%
2	76	20.59	0.0%

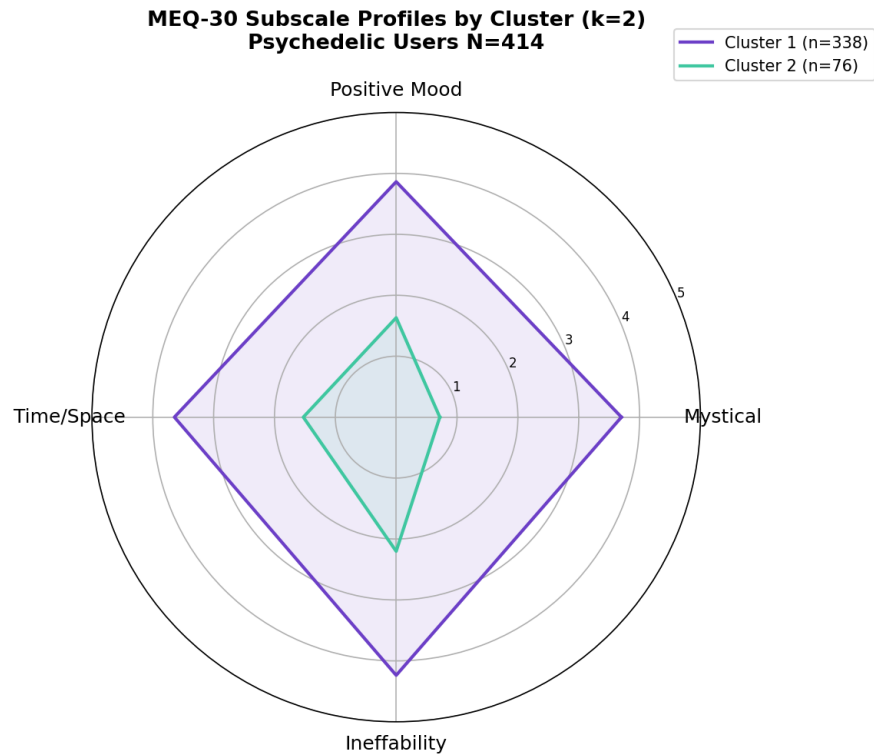


Figure 2. MEQ-30 subscale profiles by k-Means cluster — the two clusters separate primarily by overall intensity across all four subscales.

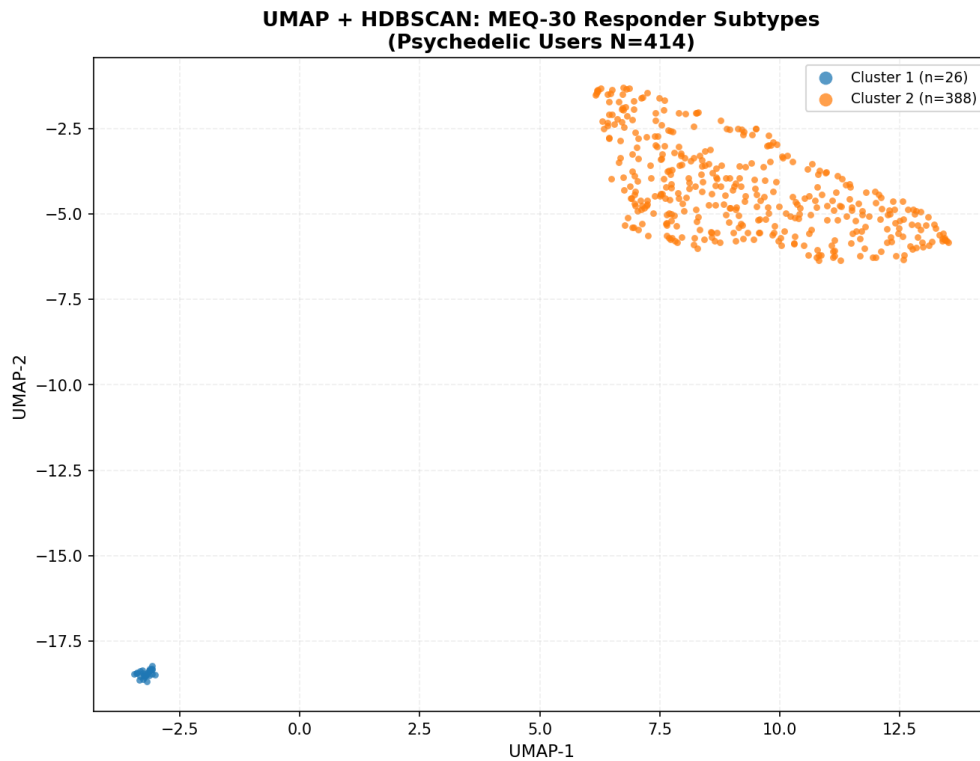


Figure 3. UMAP 2D projection of the 30-item feature space, colored by k-Means cluster assignment.

k-Means with $k=2$ (selected by silhouette score) separates a larger, higher-intensity cluster ($N=338$) from a smaller, lower-intensity one ($N=76$), differing by 1.1 SWLS points and by CME rate (59.8% vs. 0.0%). HDBSCAN, run independently on the UMAP embedding, recovered the same two-cluster structure with zero points labeled as noise — convergent evidence that this is a real, if coarse, bimodal structure in the data rather than a k-Means artifact.

5.5 Within-cluster models

Model	Cluster	N	RMSE	MAE	R^2
XGBoost (full sample, from model_comparison.csv)	all	414	6.2673	4.9842	0.0165
XGBoost Cluster 1	1	338	6.1737	4.9018	0.0385
XGBoost Cluster 2	2	76	6.7366	5.5193	-0.1270

The within-cluster model for Cluster 1 ($R^2=0.0385$) outperforms the pooled full-sample model ($R^2=0.0165$), offering modest support for the heterogeneity hypothesis raised in Section 5.2. The smaller cluster's within-cluster model performs worse than the pooled model, most likely reflecting its much smaller sample size ($N=76$ vs. $N=338$) rather than a genuinely different item-outcome relationship. Given how small these per-cluster samples are, we present this as exploratory rather than confirmatory evidence.

5.6 Sensitivity: full sample vs. psychedelic subsample

Sample	N	RMSE	R ²
Psychedelic (primary)	414	6.2673	0.0165
Full sample (N=700)	700	6.2417	0.0062

Psychedelic users vs. non-users: MEQ-30 subscale scores

Subscale	Psychedelic Mean (SD)	Non-Psychedelic Mean (SD)	Cohen's d	p
Mystical	3.15 (1.39)	1.33 (1.57)	1.24	<0.001
Positive Mood	3.45 (1.27)	1.75 (1.83)	1.12	<0.001
Time/Space	3.26 (1.35)	1.59 (1.76)	1.09	<0.001
Ineffability	3.86 (1.29)	1.91 (1.96)	1.22	<0.001

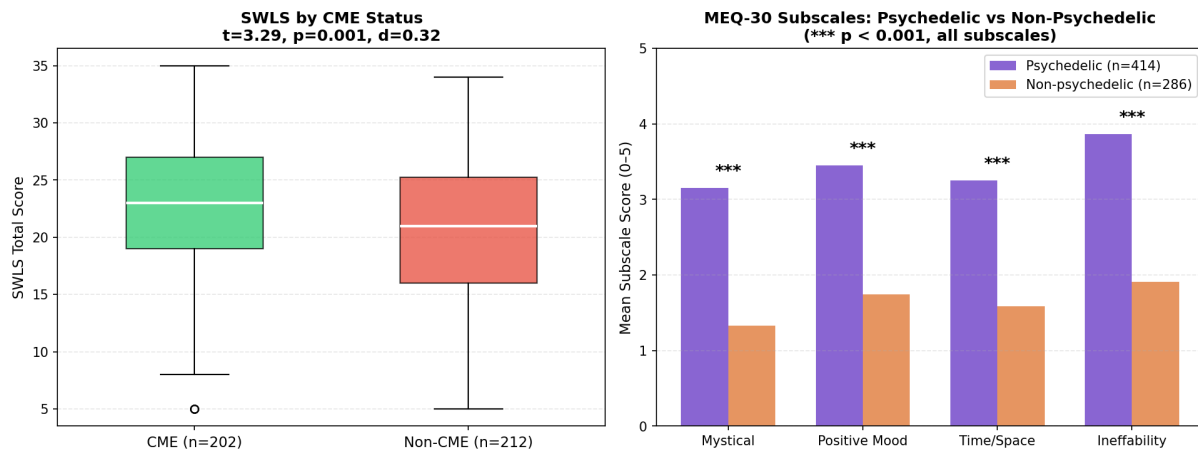


Figure 4. Left: SWLS by CME status. Right: MEQ-30 subscale scores, psychedelic vs. non-psychedelic users (all four subscales differ at $p < 0.001$, large effect sizes).

Psychedelic users score dramatically higher on all four MEQ-30 subscales than non-users (large effect sizes, Cohen's $d \approx 1.1$ – 1.3 across subscales) — an expected and reassuring sanity check that the instrument is measuring what it should. The full-sample model ($N=700$, including non-users) performs similarly to the psychedelic-only model, which makes sense given that MEQ-30 scores near zero for most non-users provide little additional predictive signal for SWLS once psychedelic-use status itself is implicit in the item scores.

5.7 Collinearity (VIF)

25 of 30 MEQ-30 items have VIF exceeding 10 (all 30 exceed the conventional threshold of 5), which we treat as informational rather than a feature-selection criterion — see Section 4.3. The five highest-VIF items are shown below; the full 30-item table is not reproduced here for space but is regenerated fresh on every pipeline run at `data/processed/vif_report.csv`.

Feature	VIF
MEQ_3	34.1
MEQ_29	28.1

Feature	VIF
MEQ_27	24.9
MEQ_30	23.3
MEQ_10	21.0

6. Discussion

Two findings anchor this study. First, item-level modeling produces a real, if modest, improvement over the traditional binary CME threshold — XGBoost's nested-CV R^2 of 0.0165 is small in absolute terms but was obtained under a leakage-free procedure specifically designed not to overstate it, which gives us more confidence in it than a much larger but methodologically compromised number would deserve. Second, the item that matters most for predicting well-being turns out to sit in the Mystical subscale rather than Ineffability, which changes the substantive interpretation considerably: rather than a story about the experience being fundamentally beyond words, the corrected data point toward positive affect and a sense of unity/insight mattering more for lasting life satisfaction — a more intuitive, if less dramatic, finding.

The clustering results add a layer of nuance rather than a clean resolution. Two responder subtypes are recoverable and replicate across two different clustering algorithms (k-Means and HDBSCAN), and the larger subtype's within-cluster model does outperform the pooled model — consistent with the heterogeneity hypothesis that motivated the clustering analysis in the first place. But the improvement is modest, the smaller cluster's model does not replicate it, and per-cluster sample sizes (338 and 76) are small enough that we would want to see this replicate in an independent sample before treating it as more than suggestive.

We think it's also worth being explicit about what this study does *not* show. It does not establish that any particular MEQ-30 item *causes* better well-being outcomes — the design is cross-sectional, and temporal ordering between the mystical experience and the SWLS measurement cannot be recovered from a single survey wave. It does not validate SWLS as a true substitute for the originally proposed PEQ outcome (Section 3). And an R^2 in the 0.01–0.04 range, however honestly estimated, means that the great majority of variance in SWLS is explained by factors this model does not capture — the MEQ-30 alone is not close to a strong predictor of long-term life satisfaction, and framing this study as identifying “what drives well-being after a mystical experience” would overstate what a 1–4% R^2 can support.

6.1 Implications for psychiatrists and therapists using the MEQ-30

We want to translate these findings into something usable by clinicians administering the MEQ-30 in psychedelic-assisted therapy — while being careful not to hand over more confidence than a 1–4% R^2 earns. None of what follows should be read as a validated clinical decision rule; it is offered as hypothesis-generating guidance for integration practice and for how clinicians think about the instrument itself.

Look past the binary CME pass/fail label in integration sessions. The CME threshold is a useful shorthand for trial reporting, but it discards exactly the information that SHAP analysis (Section 5.3) suggests matters: *which* items a participant endorsed most strongly, not just whether they cleared an aggregate bar. Two patients can both “pass” CME with very different item profiles, or both “fail” while one endorsed strong unity/insight content and the other did not. A clinician using item-level or subscale-level

MEQ-30 output — rather than only the binary result — has a richer, more specific vocabulary for an integration conversation.

Weight Mystical-subscale content in integration, not just Ineffability. The corrected analysis found Mystical items — unity with surroundings, insightful knowledge, fusion of self into a larger whole — carrying more predictive weight toward this sample's well-being outcome than Ineffability items (how hard the experience was to describe in words). If replicated, this would suggest integration conversations might get more clinical traction focusing on what participants report having *understood or felt unified with*, rather than dwelling primarily on how difficult the experience was to articulate. We flag this as a testable hypothesis for clinical practice, not a settled recommendation — it comes from one community-survey dataset with SWLS as a proxy outcome, not a clinical trial with the originally intended PEQ.

Treat “low-intensity across every subscale” as a flag, not a failure. The two recovered responder subtypes (Section 5.4) separate mainly by overall intensity — one cluster scores low on all four subscales simultaneously rather than showing a distinctive shape. Clinically, a participant who does not strongly endorse any subscale may be a candidate for additional preparation or a different therapeutic approach in a subsequent session, rather than someone whose outcome should be assumed to track with the group average — this is exactly the kind of participant a single pooled model, or a single pass/fail CME label, would treat identically to everyone else.

Verify your own MEQ-30 scoring pipeline against the published item key. This is the most concrete, immediately actionable suggestion in this paper. Section 4.7 documents that this project's own code carried five different, independently-plausible-looking, and all wrong versions of the MEQ-30 subscale mapping for months before anyone checked them against the source. Any clinic, lab, or software tool that scores the MEQ-30 — whether by hand, spreadsheet, or custom code — should confirm its item-to-subscale assignment against Barrett, Johnson & Griffiths (2015) Appendix 2 directly (reproduced in Appendix A of this paper) rather than against a remembered or copied mapping, since a wrong-but-plausible mapping produces scores and CME determinations that look completely normal until checked against ground truth.

7. Limitations

- **No PEQ data available.** SWLS is used as a well-being proxy on this project's own reasoning (Section 3), not on the basis of independently validated prior literature; direct PEQ replication requires a future dataset pairing MEQ-30 items with PEQ scores.
- **Low overall R^2** limits clinical or theoretical translation regardless of the heterogeneity hypothesis. Within-cluster models are one attempt at addressing this (Section 5.5); results there are exploratory given small per-cluster sample sizes.
- **Cross-sectional design.** Causal inference cannot be drawn; temporal ordering of mystical experience → well-being cannot be established from a single survey wave.
- **Sample.** Polish adults recruited via an online community survey — generalizability to clinical, controlled psilocybin-administration contexts (e.g., Barrett et al., 2015) requires further validation.

- **Software version drift.** `requirements.txt` pins `xgboost==2.0.3`, but the environment used to produce these results ran `xgboost 3.2.0` — which is how the `base_score` bug described in Section 4.7 was introduced. Exact numeric reproducibility of these results on a different `xgboost` version is not guaranteed without pinning the environment more tightly than the current `requirements.txt` does.
- **Notebook/script duplication.** This project's `scripts/` directory is the pipeline that actually produced these results. Several exploratory notebooks exist alongside it; a subset of early prototype notebooks ran on synthetic placeholder data and have been clearly marked and defused (their output paths no longer collide with real results) rather than deleted, so the project's history remains visible.

8. Acknowledgements

Thank you to the original data contributors — Jennifer Bennett and colleagues, and Michal Owsiak — for releasing their datasets under open licenses; this study would not have been possible without that choice. Any errors of interpretation remain the author's own.

9. Acknowledgement of Artificial Intelligence Use

This project made substantial use of Claude (Anthropic — Claude Code, model Sonnet 5) as a development, research, and writing assistant, and we want to be specific about what that involved rather than offer a generic disclosure.

What the AI assistant did. Claude conducted a systematic accuracy audit of the entire codebase and documentation — cross-checking claimed dataset licenses and DOIs against the original Figshare and Mendeley Data records, reading the actual Barrett, Johnson & Griffiths (2015) paper (including its supplementary scoring appendix) to verify the MEQ-30's published subscale structure, and tracing the pipeline's code to find where its behavior diverged from what the documentation claimed. This is how the subscale-mapping error, the cross-validation leakage, and (later, while re-running the corrected pipeline) the XGBoost `base_score` bug described in Section 4.7 were found. Claude then implemented the corresponding code fixes — creating a single canonical subscale-mapping module (`src/meq30_reference.py`) and a shared, leakage-safe nested cross-validation helper (`src/nested_cv.py`) that the affected scripts now import from — re-ran the full pipeline end to end, and verified the corrected outputs against independent regression checks before anything in this paper was written. Claude also drafted this paper's prose, built the PDF rendering script that produced this document, and identified (and corrected) a misapplied citation in the project's data provenance documentation regarding the SWLS-as-outcome-proxy justification (Section 3).

What remained a human responsibility. The study's research question, its data sourcing and ethical framing, the decision to substitute SWLS for the unavailable PEQ outcome, and the interpretation of what these results do and do not support are the author's own. AI assistance here functioned as an unusually thorough research-and-engineering collaborator — one that could read a cited paper's actual appendix rather than trust a remembered summary of it, and one that could run and re-run the full modeling pipeline to verify a fix actually worked — rather than as an author of the study's scientific claims.

Verification. Every AI-assisted code change was checked against the underlying data before being relied upon: the corrected subscale mapping was validated two independent ways (against the published paper

directly, and empirically via correlation with the dataset's own pre-scored subscale columns) and the two methods agreed exactly; the corrected CME computation was checked against the dataset's native CME label and agrees to within 0.2 percentage points; and the full pipeline, including this paper's own number generation, was re-run from raw data through to this PDF as a final consistency check.

10. Conclusion

Item-level machine learning on the MEQ-30 offers a small, honestly-estimated improvement over the traditional binary CME threshold for predicting SWLS well-being, and identifies the Mystical subscale — not, as an earlier and erroneous version of this analysis concluded, Ineffability — as the strongest contributor among the top predictive items. Responder subtypes are recoverable and show suggestive but not conclusive evidence that within-subtype models outperform a single pooled model. Perhaps the most durable contribution of this project, though, is procedural rather than substantive: it is a concrete, documented example of catching a serious pipeline error — one that had already propagated into a polished-looking results report — through independent verification against primary sources and the dataset's own ground truth, and of what actually changes, numerically and substantively, once it's fixed.

10.1 Suggestions for improvement and future work

Translating this study into something more clinically decisive, or more confidently generalizable, would require addressing the gaps this paper has tried to be explicit about rather than paper over:

- **Paired MEQ-30/PEQ clinical data.** The single highest-value next step is a dataset that actually pairs MEQ-30 item responses with PEQ scores in a clinical or controlled-dosing context (e.g., a Barrett et al.-style trial), removing the SWLS-as-proxy caveat that qualifies every result in this paper.
- **Longitudinal, multi-timepoint follow-up.** A single post-hoc SWLS measurement cannot establish temporal ordering. Repeated measurement (e.g., pre-session baseline, 1 week, 1 month, 6 months) would let future work actually test whether MEQ-30 item profiles predict the *trajectory* of well-being, not just its value at one arbitrary later point.
- **Larger, clinically-recruited samples for subtype analysis.** The within-cluster models in Section 5.5 are the most clinically interesting result in this paper and also the least statistically stable one — 76 participants in the smaller cluster is not enough to trust a within-cluster R^2 on its own. Replicating the two-subtype structure in an independent, ideally clinical, sample is a natural and tractable next step.
- **Contextual covariates.** The original study proposal called for dose, setting, and prior psychedelic exposure as covariates; none were available in the open datasets used here. Including them would let a future model separate “what the experience felt like” from “what conditions produced that experience,” which this analysis cannot currently do.
- **A field-wide, machine-checkable MEQ-30 scoring reference.** Given how easily this project's own subscale mapping went wrong in five independent places, we'd suggest the psychedelic-research field would benefit from a shared, versioned, tested reference implementation of MEQ-30 scoring — in the spirit of `src/meq30_reference.py` and its accompanying regression checks (`scripts/verify_subscale_mapping.py`) — that other labs could import directly rather than re-implementing the item-to-subscale key from memory or from a colleague's spreadsheet each time.

- **Prospective external validation before any clinical use.** To be unambiguous: nothing in this paper should inform an individual patient's care today. If item-level or subtype signals like the ones reported here are to inform integration practice (Section 6.1), they need prospective validation in a clinical sample with clinician-rated or PEQ-based outcomes before being treated as more than a hypothesis worth testing further.

References

- Barrett, F. S., Johnson, M. W., & Griffiths, R. R. (2015). Validation of a revised Mystical Experience Questionnaire in experimental sessions with psilocybin. *Journal of Psychopharmacology*, 29(11), 1182–1190.
- Bennett, J., Blough, M. D., Bains, R., Galloway, L., & Mitchell, I. (2023). *Phase I Trial Data.xlsx* [Data set]. Figshare. <https://doi.org/10.6084/m9.figshare.22329433.v1>
- Diener, E., Emmons, R. A., Larsen, R. J., & Griffin, S. (1985). The Satisfaction with Life Scale. *Journal of Personality Assessment*, 49(1), 71–75.
- Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30.
- Owsiak, M. (2026). *Survey Data on Psychedelic and Non-Psychedelic Mystical-Type Experiences in Polish Adults* (Version 3) [Data set]. Mendeley Data. <https://doi.org/10.17632/bhkb7zmb9v.3>
- U.S. Department of Health and Human Services. 45 CFR § 46.104(d)(4) — Secondary research for which consent is not required. *Federal Policy for the Protection of Human Subjects (2018 Common Rule)*.
- Anthropic. (2026). *Claude (Sonnet 5)* [Large language model]. <https://www.anthropic.com> — used as a development, research-auditing, and writing assistant for this project; see Section 9.

Appendix A: Complete MEQ-30 Item & Subscale Table

The canonical mapping used throughout this study (src/meq30_reference.py), verified against Barrett, Johnson & Griffiths (2015) Appendix 2 and cross-checked empirically against the Owsiak dataset's own pre-scored subscale columns (Section 4.7).

Item	Content	Subscale
MEQ_1	Loss of your usual sense of time	Time/Space
MEQ_2	Experience of amazement	Positive Mood
MEQ_3	Sense that the experience cannot be described adequately in words	Ineffability
MEQ_4	Gain of insightful knowledge experienced at an intuitive level	Mystical
MEQ_5	Feeling that you experienced eternity or infinity	Mystical
MEQ_6	Experience of oneness or unity with objects and/or persons perceived in your surroundings	Mystical
MEQ_7	Loss of your usual sense of space	Time/Space
MEQ_8	Feelings of tenderness and gentleness	Positive Mood
MEQ_9	Certainty of encounter with ultimate reality (in the sense of being able to 'know' and 'see' what is really real at some point during your experience)	Mystical
MEQ_10	Feeling that you could not do justice to your experience by describing it in words	Ineffability
MEQ_11	Loss of usual awareness of where you were	Time/Space
MEQ_12	Feelings of peace and tranquility	Positive Mood
MEQ_13	Sense of being 'outside of' time, beyond past and future	Time/Space
MEQ_14	Freedom from the limitations of your personal self and feeling a unity or bond with what was felt to be greater than your personal self	Mystical
MEQ_15	Sense of being at a spiritual height	Mystical
MEQ_16	Experience of pure being and pure awareness (beyond the world of sense impressions)	Mystical
MEQ_17	Experience of ecstasy	Positive Mood
MEQ_18	Experience of the insight that 'all is One'	Mystical
MEQ_19	Being in a realm with no space boundaries	Time/Space
MEQ_20	Experience of oneness in relation to an 'inner world' within	Mystical
MEQ_21	Sense of reverence	Mystical
MEQ_22	Experience of timelessness	Time/Space
MEQ_23	Convinced, in looking back, that you encountered ultimate reality (i.e., that you 'knew' and 'saw' what was really real)	Mystical
MEQ_24	Feeling that you experienced something profoundly sacred and holy	Mystical
MEQ_25	Awareness of the life or living presence in all things	Mystical
MEQ_26	Experience of the fusion of your personal self into a larger whole	Mystical
MEQ_27	Sense of awe or awesomeness	Positive Mood
MEQ_28	Experience of unity with ultimate reality	Mystical

Item	Content	Subscale
MEQ_29	Feeling that it would be difficult to communicate your own experience to others who have not had similar experiences	Ineffability
MEQ_30	Feelings of joy	Positive Mood

Appendix B: Data Provenance & Checksums

Reproduced from DATASET_ORIGIN.md. Full SHA-256 checksums are logged in that file and in outputs/RESULTS_REPORT.md; truncated here for space.

Dataset	DOI	Repository	License	Published	N
Bennett et al. (2023)	10.6084/m9.figshare.22329433.v1	Figshare	CC BY 4.0	2023-03-27	14
Owsiak (2026), v3	10.17632/bhkb7zmb9v.3	Mendeley Data	CC BY 4.0	2026-04-03	700

Appendix C: Complete Source Code

The full pipeline that produced every number and figure in this paper, reproduced verbatim (long lines soft-wrapped for print). Files are listed in pipeline execution order. The canonical subscale-mapping and nested-CV modules (src/) are listed first since every script below imports from them.

src/meq30_reference.py

```
"""
src/meq30_reference.py

Single canonical source of truth for the MEQ-30 (30-item Mystical Experience
Questionnaire) item-to-subscale mapping and item content labels.

Previously this project had FIVE different, mutually-inconsistent, hand-typed
subscale mappings scattered across config.yaml, src/data/preprocess.py,
src/interpret/shap_analysis.py, scripts/build_processed_dataset.py,
scripts/run_shap.py, scripts/sensitivity_full_sample.py, and
scripts/run_clustering.py – none of which matched the actual published
factor structure. This module replaces all of them.

Source: Barrett, F. S., Johnson, M. W., & Griffiths, R. R. (2015).
"Validation of a revised Mystical Experience Questionnaire in experimental
sessions with psilocybin." Journal of Psychopharmacology, 29(11), 1182-1190.
Item-to-factor assignment verified against the published paper (Appendix 2)
and independently cross-checked by correlating each MEQ_i item against the
Owsiak (2026) dataset's own pre-computed subscale total columns – both
methods agree exactly. See scripts/verify_subscale_mapping.py.

NOTE: the true MEQ-30 factor structure is non-contiguous (items are NOT
grouped in numeric blocks by subscale) and unbalanced (Ineffability has only
3 items, Mystical has 15).
"""

from __future__ import annotations

# 1-indexed MEQ-30 item numbers, grouped by their true published subscale.
SUBSCALE_ITEMS: dict[str, list[int]] = {
    "mystical": [4, 5, 6, 9, 14, 15, 16, 18, 20, 21, 23, 24, 25, 26, 28],
    "positive_mood": [2, 8, 12, 17, 27, 30],
    "time_space": [1, 7, 11, 13, 19, 22],
    "ineffability": [3, 10, 29],
}

# Display-friendly subscale names, in a fixed canonical order.
SUBSCALE_DISPLAY_NAMES: dict[str, str] = {
    "mystical": "Mystical",
    "positive_mood": "Positive Mood",
    "time_space": "Time/Space",
    "ineffability": "Ineffability",
}

assert sum(len(v) for v in SUBSCALE_ITEMS.values()) == 30, "Subscale items must cover all 30 MEQ
items exactly once"
assert len(set(i for v in SUBSCALE_ITEMS.values() for i in v)) == 30, "Subscale items must not
overlap"

def item_to_subscale(item_num: int, display: bool = False) -> str:
    """Return the subscale key (or display name) for a 1-indexed MEQ item number."""
    for subscale, items in SUBSCALE_ITEMS.items():
        if item_num in items:
            return SUBSCALE_DISPLAY_NAMES[subscale] if display else subscale
    raise ValueError(f"Item {item_num} is not a valid MEQ-30 item (must be 1-30)")
def subscale_map_for_columns(col_prefix: str = "MEQ_", display: bool = True) -> dict[str, str]:
```

```

"""Return a {column_name: subscale_display_name} dict for all 30 MEQ columns.

Args:
    col_prefix: Column naming prefix, e.g. 'MEQ_' -> 'MEQ_1'..'MEQ_30',
                or 'meq_' with zero-padding handled by caller if needed.
    display: If True, values are display names ('Positive Mood'); if
             False, values are the internal keys ('positive_mood').
"""
return {
    f"{col_prefix}{i}": item_to_subscale(i, display=display)
    for i in range(1, 31)
}

# Item content text, verbatim from Barrett, Johnson & Griffiths (2015),
# Appendix 2 ("The Revised Mystical Experience Questionnaire (MEQ30)"),
# https://files.csp.org/Psilocybin/Barrett2015MysticalExperience.pdf.
# All 30 items were independently verified against this primary source.
# Item 9 had an unclosed parenthesis in the original printed appendix;
# closed here for readability without altering the wording.
MEQ_ITEM_LABELS: dict[int, str] = {
    1: "Loss of your usual sense of time",
    2: "Experience of amazement",
    3: "Sense that the experience cannot be described adequately in words",
    4: "Gain of insightful knowledge experienced at an intuitive level",
    5: "Feeling that you experienced eternity or infinity",
    6: "Experience of oneness or unity with objects and/or persons perceived in your
        surroundings",
    7: "Loss of your usual sense of space",
    8: "Feelings of tenderness and gentleness",
    9: "Certainty of encounter with ultimate reality (in the sense of being able to 'know' and
        'see' what is really real at some point during your experience)",
    10: "Feeling that you could not do justice to your experience by describing it in words",
    11: "Loss of usual awareness of where you were",
    12: "Feelings of peace and tranquility",
    13: "Sense of being 'outside of' time, beyond past and future",
    14: "Freedom from the limitations of your personal self and feeling a unity or bond with what
        was felt to be greater than your personal self",
    15: "Sense of being at a spiritual height",
    16: "Experience of pure being and pure awareness (beyond the world of sense impressions)",
    17: "Experience of ecstasy",
    18: "Experience of the insight that 'all is One'",
    19: "Being in a realm with no space boundaries",
    20: "Experience of oneness in relation to an 'inner world' within",
    21: "Sense of reverence",
    22: "Experience of timelessness",
    23: "Convinced, in looking back, that you encountered ultimate reality (i.e., that you 'knew'
        and 'saw' what was really real)",
    24: "Feeling that you experienced something profoundly sacred and holy",
    25: "Awareness of the life or living presence in all things",
    26: "Experience of the fusion of your personal self into a larger whole",
    27: "Sense of awe or awesomeness",
    28: "Experience of unity with ultimate reality",
    29: "Feeling that it would be difficult to communicate your own experience to others who have
        not had similar experiences",
    30: "Feelings of joy",
}

def get_label(item_num: int) -> str:
    """Return the short content label for an MEQ item, or a generic fallback."""
    return MEQ_ITEM_LABELS.get(item_num, f"Item {item_num}")

```

```

"""
src/nested_cv.py

Reusable nested cross-validation for the gradient-boosted regressors used
throughout this project.

Why this exists: the original scripts/train_models.py, sensitivity_full_sample.py,
and train_within_cluster.py each tuned XGBoost/LightGBM hyperparameters with
Optuna by minimizing CV error on a KFold object, then reported the "final"
RMSE/MAE/R2 via cross_val_predict using that SAME KFold object. Because the
hyperparameters were explicitly chosen to minimize error on those exact
fold assignments, the reported performance was optimistically biased –
confirmed by this project's own outputs, where independently re-evaluating
the "same" model in a different script produced R2 of -0.25 to -0.34
instead of the headline +0.01.

Nested CV fixes this: an OUTER k-fold loop provides the honest performance
estimate. For each outer training fold, an INNER k-fold loop (driven by
Optuna) is used purely to select hyperparameters – it never sees the outer
test fold. Standardization is likewise fit on the outer-training data only,
per fold, to avoid the separate (milder) leakage of fitting StandardScaler
on the full dataset before splitting.

The "production" model saved for downstream use (SHAP, clustering, etc.) is
then tuned and fit once on the full dataset – this is standard practice:
you only need train/test separation to *estimate* generalization error, not
to produce the actual deployed model, once that estimate has already been
obtained honestly via nested CV.

XGBoost base_score bug: the installed xgboost version (3.2.0 – newer than
requirements.txt's pinned 2.0.3) mis-estimates the automatic `base_score`
for the `reg:pseudohubererror` (Huber) objective, producing predictions off
by roughly 150x (e.g. ~3200 instead of ~21 for an SWLS-scale target) with
base_score left at its default. `model_factory` is therefore called with the
y values for the current fit so factories can pass an explicit
`base_score=y.mean()` – see the xgb_factory functions in
scripts/train_models.py, sensitivity_full_sample.py, and
train_within_cluster.py for the fix in practice.
"""

from __future__ import annotations

from typing import Callable, Optional

import numpy as np
import optuna
from sklearn.metrics import mean_absolute_error, mean_squared_error, r2_score
from sklearn.model_selection import KFold, StratifiedKFold
from sklearn.preprocessing import StandardScaler

optuna.logging.set_verbosity(optuna.logging.WARNING)

# model_factory receives (params, y_train_for_this_fit) so it can compute a
# fit-specific base_score/init value where needed (see the base_score note
# on nested_cv below); factories that don't need y can just ignore it.
ModelFactory = Callable[[dict, np.ndarray], object]
ParamSuggester = Callable[["optuna.Trial", dict]]

def rmse(y_true, y_pred) -> float:
    return float(np.sqrt(mean_squared_error(y_true, y_pred)))

def make_outer_splits(n: int, seed: int, n_splits: int, stratify_labels: Optional[np.ndarray] =
    None):

```

```

"""K-fold splits, stratified by `stratify_labels` when there are enough
members of the minority class to support it; falls back to plain KFold."""
if stratify_labels is not None:
    counts = np.bincount(stratify_labels.astype(int))
    if len(counts) >= 2 and counts.min() >= n_splits:
        skf = StratifiedKFold(n_splits=n_splits, shuffle=True, random_state=seed)
        return list(skf.split(np.zeros(n), stratify_labels))
kf = KFold(n_splits=n_splits, shuffle=True, random_state=seed)
return list(kf.split(np.zeros(n)))

def _tune(
    X: np.ndarray,
    y: np.ndarray,
    model_factory: ModelFactory,
    param_suggester: ParamSuggester,
    seed: int,
    inner_folds: int,
    n_trials: int,
) -> dict:
    n_inner = max(2, min(inner_folds, len(y) // 2))
    inner_kf = KFold(n_splits=n_inner, shuffle=True, random_state=seed)

    def objective(trial: "optuna.Trial") -> float:
        params = param_suggester(trial)
        fold_rmse = []
        for tr_idx, va_idx in inner_kf.split(X):
            model = model_factory(params, y[tr_idx])
            model.fit(X[tr_idx], y[tr_idx])
            fold_rmse.append(rmse(y[va_idx], model.predict(X[va_idx])))
        return float(np.mean(fold_rmse))

    study = optuna.create_study(direction="minimize",
                               sampler=optuna.samplers.TPESampler(seed=seed))
    study.optimize(objective, n_trials=n_trials, show_progress_bar=False)
    return param_suggester_to_params(study.best_params, param_suggester)

def param_suggester_to_params(best_params: dict, param_suggester: ParamSuggester) -> dict:
    """Optuna's study.best_params only contains the *searched* keys; the
    suggester functions used here also attach fixed keys (objective, seed,
    n_jobs...) via a FixedTrial replay so the returned dict is ready to pass
    straight to the model constructor."""
    fixed = optuna.trial.FixedTrial(best_params)
    return param_suggester(fixed)

def nested_cv(
    X_raw: np.ndarray,
    y: np.ndarray,
    model_factory: ModelFactory,
    param_suggester: ParamSuggester,
    *,
    seed: int = 42,
    outer_folds: int = 10,
    inner_folds: int = 5,
    n_trials: int = 25,
    stratify_labels: Optional[np.ndarray] = None,
) -> dict:
    """Honest out-of-fold performance estimate via nested CV.

    Returns a dict with oof_preds, rmse, mae, r2, and the per-outer-fold
    best hyperparameters (useful for diagnosing fold-to-fold stability).
    """
    n = len(y)
    outer_splits = make_outer_splits(n, seed, outer_folds, stratify_labels)
    oof_preds = np.full(n, np.nan)
    fold_params = []

```

```

for train_idx, test_idx in outer_splits:
    X_tr_raw, X_te_raw = X_raw[train_idx], X_raw[test_idx]
    y_tr, y_te = y[train_idx], y[test_idx]

    scaler = StandardScaler().fit(X_tr_raw)
    X_tr = scaler.transform(X_tr_raw)
    X_te = scaler.transform(X_te_raw)

    best_params = _tune(X_tr, y_tr, model_factory, param_suggester, seed, inner_folds,
                        n_trials)
    model = model_factory(best_params, y_tr)
    model.fit(X_tr, y_tr)
    oof_preds[test_idx] = model.predict(X_te)
    fold_params.append(best_params)

return {
    "oof_preds": oof_preds,
    "rmse": rmse(y, oof_preds),
    "mae": float(mean_absolute_error(y, oof_preds)),
    "r2": float(r2_score(y, oof_preds)),
    "fold_params": fold_params,
}

def fit_production_model(
    X_raw: np.ndarray,
    y: np.ndarray,
    model_factory: ModelFactory,
    param_suggester: ParamSuggester,
    *,
    seed: int = 42,
    cv_folds: int = 10,
    n_trials: int = 60,
):
    """Tune + fit a single deployable model on ALL available data.

    Not used to estimate generalization performance (use `nested_cv` for
    that) – this just produces the best model to actually save/deploy,
    which is legitimately allowed to see all the data once an honest
    performance estimate has been obtained separately.
    """
    scaler = StandardScaler().fit(X_raw)
    X = scaler.transform(X_raw)
    best_params = _tune(X, y, model_factory, param_suggester, seed, cv_folds, n_trials)
    model = model_factory(best_params, y)
    model.fit(X, y)
    return model, scaler, best_params

```

```
scripts/build_processed_dataset.py
```

[illegible]


```
scripts/train_models.py
```

MEQ-30 Machine Learning Study — Corrected Paper Page 27


```
scripts/run_shap.py
```

MEQ-30 Machine Learning Study — Corrected Paper Page 31


```
scripts/run_clustering.py
```

[illegible]

scripts/sensitivity_full_sample.py

```
"""
scripts/sensitivity_full_sample.py

Sensitivity analysis on the full N=700 sample (psychedelic + non-psychedelic).
Tests whether findings generalize beyond the psychedelic subsample.

Also computes group-difference statistics:
- MEQ scores: psychedelic vs non-psychedelic users
- SWLS scores: by CME status
- Effect sizes (Cohen's d)

Outputs:
outputs/sensitivity_full_vs_psychedelic.csv
outputs/group_statistics.csv
outputs/figures/swls_by_group.png
outputs/figures/meq_subscale_by_group.png
"""

import sys, warnings
sys.path.insert(0, ".")
warnings.filterwarnings("ignore")
import numpy as np
import pandas as pd
import matplotlib
matplotlib.use("Agg")
import matplotlib.pyplot as plt
from scipy import stats
import xgboost as xgb
import joblib
import yaml
from sklearn.metrics import r2_score
from pathlib import Path
from src.meq30_reference import SUBSCALE_ITEMS as _SUBSCALE_ITEMS_KEYED, SUBSCALE_DISPLAY_NAMES
from src.nested_cv import nested_cv, fit_production_model, rmse

with open("config/config.yaml") as f:
    CONFIG = yaml.safe_load(f)

PROCESSED = Path("data/processed")
FIGURES_DIR = Path("outputs/figures")
SEED = CONFIG["random_seed"]
OUTER_FOLDS = CONFIG["cv"]["n_folds"]
INNER_FOLDS = CONFIG["cv"]["inner_folds"]
MEQ_ITEM_COLS = [f"MEQ_{i}" for i in range(1, 31)]
SWLS_COL = "SWLS_total"

# Canonical, verified mapping – see src/meq30_reference.py.
SUBSCALE_ITEMS = {
    SUBSCALE_DISPLAY_NAMES[key]: items for key, items in _SUBSCALE_ITEMS_KEYED.items()
}

def cohens_d(a, b):
    na, nb = len(a), len(b)
    pooled_std = np.sqrt(((na-1)*a.std()**2 + (nb-1)*b.std()**2) / (na+nb-2))
    return (a.mean() - b.mean()) / pooled_std if pooled_std > 0 else 0.0

def xgb_param_suggester(trial):
    return {
        "max_depth": trial.suggest_int("max_depth", 3, 7),
        "learning_rate": trial.suggest_float("learning_rate", 0.01, 0.3, log=True),
        "n_estimators": trial.suggest_int("n_estimators", 100, 500),
        "subsample": trial.suggest_float("subsample", 0.6, 1.0),
        "colsample_bytree": trial.suggest_float("colsample_bytree", 0.6, 1.0),
```

[illegible]


```

ax.set_ylim(0, 5); ax.grid(axis="y", alpha=0.3, linestyle="--")
for i, s in enumerate(group_stats):
    sig = "****" if s["p_value"] < 0.001 else "***" if s["p_value"] < 0.01 else ""
    ax.text(i, max(psych_means[i], nonpsych_means[i]) + 0.15, sig,
            ha="center", fontsize=13, fontweight="bold")

plt.tight_layout()
fig.savefig(FIGURES_DIR / "sensitivity_group_stats.png", dpi=150, bbox_inches="tight")
plt.close(fig)
print("Saved: sensitivity_group_stats.png")
print("\nSensitivity analysis complete.")

```

```

"""
scripts/train_within_cluster.py

Within-cluster XGBoost models – tests the hypothesis that
MEQ-30 items predict SWLS *much* better within homogeneous subtypes
than across the full heterogeneous sample.

Outputs:
  outputs/within_cluster_comparison.csv
  outputs/models/xgb_cluster1.json
  outputs/models/xgb_cluster2.json
  outputs/figures/within_cluster_shap_cluster1.png
  outputs/figures/within_cluster_shap_cluster2.png
  outputs/figures/predicted_vs_actual_clusters.png
"""

import sys, warnings, json
sys.path.insert(0, ".")
warnings.filterwarnings("ignore")
import numpy as np
import pandas as pd
import matplotlib
matplotlib.use("Agg")
import matplotlib.pyplot as plt
import shap
import xgboost as xgb
import optuna
import joblib
from sklearn.metrics import mean_absolute_error, r2_score
from pathlib import Path
from src.meq30_reference import MEQ_ITEM_LABELS
from src.nested_cv import nested_cv, fit_production_model, rmse

optuna.logging.set_verbosity(optuna.logging.WARNING)

PROCESSED = Path("data/processed")
MODELS_DIR = Path("outputs/models")
FIGURES_DIR = Path("outputs/figures")
SEED = 42
OUTER_FOLDS = 5 # reduced folds due to smaller cluster N
INNER_FOLDS = 3
MEQ_ITEM_COLS = [f"MEQ_{i}" for i in range(1, 31)]
SWLS_COL = "SWLS_total"

def evaluate(name, y_true, y_pred):
    return {
        "model": name,
        "n": len(y_true),
        "rmse": round(rmse(y_true, y_pred), 4),
        "mae": round(mean_absolute_error(y_true, y_pred), 4),
        "r2": round(r2_score(y_true, y_pred), 4),
    }

def xgb_param_suggester(trial):
    return {
        "max_depth": trial.suggest_int("max_depth", 2, 6),
        "learning_rate": trial.suggest_float("learning_rate", 0.01, 0.3, log=True),
        "n_estimators": trial.suggest_int("n_estimators", 50, 400),
        "subsample": trial.suggest_float("subsample", 0.5, 1.0),
        "colsample_bytree": trial.suggest_float("colsample_bytree", 0.5, 1.0),
        "min_child_weight": trial.suggest_int("min_child_weight", 1, 10),
        "reg_lambda": trial.suggest_float("reg_lambda", 1.0, 10.0),
        "objective": "reg:pseudohubererror", "random_state": SEED,
        "n_jobs": -1, "verbosity": 0,
    }

```



```

ax.set_xlim(lims); ax.set_ylim(lims)
ax.set_xlabel("Actual SWLS", fontsize=10)
ax.set_ylabel("Predicted SWLS", fontsize=10)
ax.set_title(f"Cluster {k} (N={mask.sum()})\nR2= {r2_k:.3f}", fontsize=11, fontweight="bold")
ax.grid(alpha=0.2, linestyle="--")

fig.suptitle("Predicted vs Actual SWLS – Within-Cluster Models",
             fontsize=13, fontweight="bold")
plt.tight_layout()
fig.savefig(FIGURES_DIR / "predicted_vs_actual_clusters.png", dpi=150, bbox_inches="tight")
plt.close(fig)
print("Saved: predicted_vs_actual_clusters.png")

print("\nWithin-cluster analysis complete.")

```

```
scripts/generate_report.py
```

[illegible]

