

Machine learning for communicating water risk and predicting public quality categories

Aprendizaje automático para comunicar el riesgo hídrico y predecir las categorías públicas de calidad

Recibido: 20/04/2026 - Aceptado: 10/07/2026

Pedro Alvaro Edwin Gallegos Pasco

<https://orcid.org/0000-0002-3563-1991>

paegallegos@unap.edu.pe

Universidad Nacional del Altiplano. Puno, Perú

Haydee Celia Pineda Chaíña

<https://orcid.org/0000-0002-9112-9277>

hcpineda@unap.edu.pe

Universidad Nacional del Altiplano. Puno, Perú

Luz Marina Teves Ponce

<https://orcid.org/0009-0002-6130-7092>

ltevesp@unap.edu.pe

Universidad Nacional del Altiplano. Puno, Perú

Alfredo Mamani Canqui

<https://orcid.org/0000-0002-8379-8366>

amamani@unap.edu.pe

Universidad Nacional del Altiplano. Puno, Perú

Nestor Omar Mercado Ayamamani

<https://orcid.org/0000-0002-9120-939X>

Nestor@unap.edu.pe

Universidad Nacional del Altiplano. Puno, Perú

Johan Denis Aguilar Ramirez

<https://orcid.org/0009-0001-5371-6526>

Jaguilarr@est.unap.edu.pe

Universidad Nacional del Altiplano. Puno, Perú

Abstract

Water-quality indices are not only technical summaries but communication devices: they translate complex multivariate monitoring data into a few public categories used to understand and act upon environmental risk. This article reinterprets the Canadian Council of Ministers of the Environment Water Quality Index (CCME-WQI) as a communication instrument and asks how faithfully its public categories can be reconstructed from raw physicochemical data, and what is lost when a continuous risk gradient is compressed into a public message. Using 400 observations from 12 stations of Peru's National Water Authority in the Llave River Basin, supervised classifiers and a stacking ensemble were trained to recover the three public categories (Excellent, Good, Fair/Poor), while silhouette analysis and principal-component analysis probed the data structure. The categories proved only moderately recoverable: the best model reached 65% exact accuracy, a ceiling no method surpassed, yet about 91% within-one-category accuracy, almost never confusing excellent with degraded water. Nearly half the observations sit at the index ceiling and most of the rest sit just below a threshold, so distinct public messages attach to near-identical water; a near-zero category silhouette and PCA confirm a high-dimensional continuum that

the labels cut across. We argue that this gap between gradient and label is the price of legibility, and discuss how category boundaries frame public risk perception and the unequal capacity of Quechua- and Aymara-speaking communities to access these messages.

Keywords: environmental risk communication, water quality index, machine learning, science communication.

Resumen

Los índices de calidad del agua no son solo resúmenes técnicos, sino también herramientas de comunicación: traducen datos complejos de monitoreo multivariado en unas pocas categorías públicas que se utilizan para comprender y actuar ante el riesgo ambiental. Este artículo reinterpreta el Índice de Calidad del Agua del Consejo Canadiense de Ministros del Medio Ambiente (CCME-WQI) como un instrumento de comunicación y se pregunta con qué fidelidad se pueden reconstruir sus categorías públicas a partir de datos fisicoquímicos brutos, y qué se pierde cuando un gradiente de riesgo continuo se comprime en un mensaje público. Utilizando 400 observaciones de 12 estaciones de la Autoridad Nacional del Agua de Perú en la Cuenca del Río Llave, se entrenaron clasificadores supervisados y un conjunto de apilamiento para recuperar las tres categorías públicas (Excelente, Buena, Regular/Mala), mientras que el análisis de silueta y el análisis de componentes principales exploraron la estructura de los datos. Las categorías demostraron ser solo moderadamente recuperables: el mejor modelo alcanzó una precisión exacta del 65%, un límite que ningún método superó, pero una precisión dentro de una categoría de aproximadamente el 91%, casi nunca confundiendo agua excelente con agua degradada. Casi la mitad de las observaciones se sitúan en el límite superior del índice y la mayoría del resto justo por debajo de un umbral, lo que indica que distintos mensajes públicos se asocian a situaciones casi idénticas. Una silueta de categoría cercana a cero y un análisis de componentes principales (PCA) confirman un continuo de alta dimensionalidad que atraviesan las etiquetas. Sostenemos que esta brecha entre el gradiente y la etiqueta es el precio de la legibilidad, y analizamos cómo los límites de las categorías configuran la percepción pública del riesgo y la capacidad desigual de las comunidades quechua y aimaras para acceder a estos mensajes.

Palabras clave: comunicación del riesgo ambiental, índice de calidad del agua, aprendizaje automático, comunicación pública de la ciencia.

Introduction

Environmental risk is, above all, a problem of communication. Modern societies generate vast quantities of technical data about the hazards that surround them, yet those data become socially meaningful only once they are translated into messages that publics, decision-makers and affected communities can interpret and act upon (Beck, 1992). The governance of environmental hazards thus depends not merely on measurement, but on the communicative work of rendering measurement intelligible (Renn, 2008). Few domains illustrate this as sharply as water quality, where dozens of physicochemical parameters — many interacting in nonlinear ways — must be condensed into a form that non-specialists can grasp.

Composite indices perform exactly this translation. By collapsing a multivariate measurement space into a single score and a small set of labelled categories, indices such as the Canadian Council of Ministers of the Environment Water Quality Index (CCME-WQI) make the otherwise invisible state of a water body publicly legible (Banda & Kumarasamy, 2020). Seen this way, an index is less a neutral instrument than a message: it selects which aspects of reality to foreground, names them, and frames how audiences perceive a situation (Entman, 1993). It also operates as a boundary object that lets scientists, regulators and citizens coordinate around a shared — if simplified — representation of the environment (Star & Griesemer, 1989).

How scientific information reaches the public has long preoccupied science-communication scholarship. Early deficit models assumed that supplying more facts would close the gap between experts and lay audiences; later work showed, instead, that publics actively reinterpret information through their own knowledge, identities and circumstances (Wynne, 1992). Dialogic and participatory approaches therefore came to stress the co-production of environmental knowledge and the value of local and citizen expertise (Irwin, 1995), reframing the public communication of science as a relational process rather than a one-way transfer (Bucchi & Trench, 2014).

Environmental communication research has documented how media and institutional framings shape collective understanding of ecological issues, from climate change to resource extraction (Massa et al., 2024). The

frames chosen and the language used to describe a hazard influence not only perception but the political salience of a problem (Chen et al., 2024; Sandoval et al., 2025). The social amplification of risk framework, in turn, shows that risk signals are intensified or muted as they pass through informational and institutional channels, so that the public impact of a hazard frequently diverges from its technical magnitude (Kasperson et al., 1988). Effective risk communication is, consequently, a sustained effort to build shared understanding under uncertainty not a single act of disclosure (Fischhoff, 1995).

These concerns carry particular weight in the Peruvian Altiplano. The Ilave River Basin, in the Puno region, is one of the four main Peruvian tributaries of Lake Titicaca and supplies water for irrigation, domestic use and livestock across territories inhabited largely by Quechua- and Aymara-speaking communities. Water quality there is shaped both by natural geochemistry — the weathering of volcanic and sulfide-rich minerals that raises baseline arsenic and heavy-metal levels — and by anthropogenic loads from wastewater, agriculture and small-scale mining (Biamont-Rojas et al., 2023; Salas-Ávila et al., 2021). Health-risk assessments across southern Peruvian basins have repeatedly flagged dissolved metals as a public concern (Escobar-Mamani et al., 2023). Monitoring follows standardized national protocols (Autoridad Nacional del Agua, 2016) and is evaluated against national environmental quality standards (Ministerio del Ambiente, 2017); yet the resulting information seldom reaches affected communities in an accessible form, raising a question of information justice — who can obtain, understand and use environmental data, and in which language.

Two bodies of work bear on this problem without quite meeting. On one side, machine learning has become a powerful means of classifying water quality from multivariate data, with ensemble methods frequently outperforming single-algorithm classifiers (Abbas et al., 2024; Frincu, 2025; Zhu et al., 2022) and with supervised and unsupervised techniques combined to recover both labelled categories and latent structure (Krishnaraj & Deka, 2020; Nasir et al., 2022). On the other, the index literature treats the CCME-WQI chiefly as a technical artefact, debating its formulation and thresholds (Mogane et al., 2023; Swain, 2024). Neither has examined the index as a communication device — how classifying for public consumption frames risk, what it discards, and what this means socially. To our knowledge, no study has approached water-quality classification in a high-altitude Andean watershed from a communication-and-society perspective.

This article addresses that gap by reframing automated water-quality classification as an instance of environmental-risk communication. Three questions guide the analysis. First, how faithfully can the public categories of the CCME-WQI be reconstructed from raw physicochemical data using machine learning? Second, what does the (mis)alignment between the natural data structure and the index categories reveal about the simplification inherent in communicating risk? Third, what are the social and communicative implications for the Andean communities that depend on these messages? The contribution is threefold: a conceptual bridge between data-driven indices and communication scholarship; empirical evidence on the fidelity and limits of the categories that publics actually receive; and design and governance implications for communicating water risk in the Altiplano.

Water-quality indices as communication artefacts

Treating an index as a message reframes the analytical questions. A water-quality index does three communicative things at once. First, it selects: from fourteen or more parameters it foregrounds those that breach guidelines and leaves the rest in the background, an operation of salience that lies at the heart of framing (Entman, 1993). Second, it categorizes: continuous scores are cut into named classes — Excellent, Good, Fair/Poor — that become the actual currency of public communication, since few non-specialists engage with the underlying numbers. Third, it coordinates: by offering a common label, the index allows scientists, the National Water Authority (ANA), municipal authorities and water users to align their understanding, working as a boundary object that travels across these communities while meaning something slightly different to each (Star & Griesemer, 1989).

Each of these operations entails loss. Categorization draws sharp boundaries through a continuous gradient, so that two observations on either side of a threshold may be communicated very differently despite being almost identical. This is not a flaw to be engineered away, but the price of legibility: simplification is what makes the message communicable at all. The pertinent question for risk communication is whether that simplification amplifies or attenuates the signal relative to the underlying hazard (Kasperson et al., 1988), and whether the resulting message reaches and is understood by those who most need it (Fischhoff, 1995; Wynne, 1992). The analysis that follows makes the trade-off visible: supervised models gauge how recoverable the public categories are from the data, while unsupervised methods gauge how closely those categories track the structure of the data themselves.

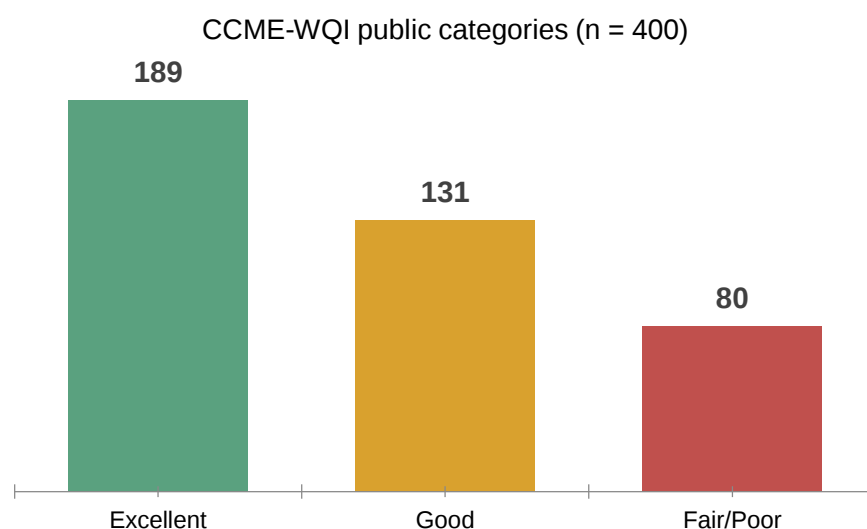
Methodology

Study area and data

The study draws on physicochemical measurements collected by the ANA's Subdirectorate of Water Resources Quality Management between 2015 and 2019 at 12 stations along the main channel and tributaries of the llave River, in line with the National Monitoring Protocol (Autoridad Nacional del Agua, 2016). The dataset comprises 400 observations and 14 variables: phosphorus, nine metals and metalloids (As, Cr, Cd, Cu, Fe, Mn, Hg, Pb, Zn), total coliforms, biochemical oxygen demand, pH and electrical conductivity sampled across wet-season ($n = 222$) and dry-season ($n = 178$) campaigns. For each observation the CCME-WQI was computed from the standard three-factor formulation of scope, frequency and amplitude (Banda & Kumarasamy, 2020) and grouped into three public categories: Excellent (≥ 95), Good (80–94.9) and Fair/Poor (< 80), the last merging the original Fair, Marginal and Poor labels for analytical stability (Mogane et al., 2023). The resulting distribution was 189 Excellent (47.2%), 131 Good (32.8%) and 80 Fair/Poor (20.0%) (Figure 1).

Figure 1

Distribution of CCME-WQI public categories



Reconstructing the public categories

Recovering the index's public categories was framed as a multi-class classification problem mapping the 14 features onto the three labels. To prevent leakage, the CCME-WQI score itself was excluded from the predictors (Abbas et al., 2024). Strongly right-skewed concentration variables (metals, total coliforms and BOD) were log-transformed and all features standardized, with every preprocessing step fitted inside the cross-validation folds to avoid leakage; because Fair/Poor is the smallest class, the training folds were rebalanced with SMOTE (Synthetic Minority Over-sampling Technique; Chawla et al., 2002). Five classifiers were compared — Random Forest (Biau & Scornet, 2016), histogram-based gradient boosting (Bentéjac et al., 2021), extreme gradient boosting (XGBoost; Chen & Guestrin, 2016) and a Support Vector Machine (Cervantes et al., 2020), together with a stacking ensemble that combines them through a logistic meta-learner (Wolpert, 1992), all implemented in Python with scikit-learn (Scikit-learn developers, 2024) and tuned with stratified cross-validation. Performance was assessed by 10-fold stratified cross-validation, reporting exact accuracy, macro-averaged F1 and the recall of the minority Fair/Poor class, which carries the most consequential public warning. Because the three categories are ordered (Excellent > Good > Fair/Poor), we additionally report within-one-category (adjacent) accuracy, which counts a prediction as correct when it falls on the true category or an immediately neighbouring one, separating minor boundary errors from the rare and serious confusion of excellent with degraded water.

Probing the underlying structure

To examine whether the public categories correspond to natural structure in the data, we computed the silhouette coefficient of the three CCME-WQI categories in the standardized feature space — a direct measure of how separable the categories are, where values near or below zero indicate that they are not well-defined groups. Dimensionality was examined with Principal Component Analysis (PCA), used to assess how variance is distributed across orthogonal components (Krishnaraj & Deka, 2020). Together, the silhouette of the categories and the dimensional structure revealed by PCA offer a direct, empirical measure of how much the communicative categories simplify the environmental gradient.

Software and artificial-intelligence disclosure

All analyses were carried out in Python using scikit-learn, XGBoost, LightGBM, CatBoost and imbalanced-learn, and the figures were produced with Matplotlib. The water-quality measurements are real monitoring data from Peru's National Water Authority (ANA); no data were generated, synthesised or imputed by artificial intelligence. In the interest of transparency, the authors declare that generative artificial-intelligence tools were used as a support during the study: to assist in writing and revising the analysis code, in producing the figures, and in editing the language of the manuscript. The authors designed the research, defined all analytical choices, and verified, interpreted and assume full responsibility for every result reported here; artificial intelligence is not, and could not be, an author of this work. The complete dataset is provided as supplementary material so that the analysis can be fully reproduced.

Results and discussion

How recoverable is the public message?

Under 10-fold cross-validation the public categories proved only moderately recoverable, and no method broke past that limit: a stacking ensemble of tree- and boosting-based learners reached the highest exact accuracy (65.0%; macro-F1 = 0.623), marginally ahead of extreme gradient boosting (63.5%), histogram-based gradient boosting (63.2%) and Random Forest (62.7%), with the Support Vector Machine trailing (60.0%) (Table 1). That even a carefully tuned ensemble, with class rebalancing and skew correction, recovers the exact category barely two times in three is itself the finding: the labels track real structure in the measurements but do not sit cleanly within it, which already signals that the public message simplifies more than it reveals (Swain, 2024). Yet the picture shifts once the ordinal nature of the categories is taken into account: measured as within-one-category (adjacent) accuracy, every model exceeds 89% and the best reach 91.2%, so the classifiers almost never commit the one truly serious error labelling excellent water as degraded, or the reverse. The public message, in short, is hard to reproduce exactly but easy to reproduce approximately, which is precisely what one expects when sharp labels are laid over a continuous gradient.

Table 1

Reconstruction of the public categories

Model	Exact accuracy	Within-1 accuracy	Macro-F1	Fair/Poor recall
Random Forest	0.627 ± 0.077	0.908	0.604	0.550
Histogram gradient boosting	0.632 ± 0.073	0.910	0.604	0.500
XGBoost	0.635 ± 0.071	0.912	0.610	0.525
SVM	0.600 ± 0.057	0.895	0.577	0.525
Stacking ensemble	0.650 ± 0.072	0.912	0.623	0.512

The minority Fair/Poor class remained the hardest to recover even after rebalancing. Across the cross-validated models its recall settled around one half best for the balanced Random Forest at 0.55 — so roughly half of the most degraded samples were still assigned to a milder category. Because Fair/Poor is precisely the message that warns of the most degraded water, a model's sensitivity to this class is, in communicative terms, its capacity to

keep the alarm from falling silent, and these residual misses attenuate the risk signal reaching the public, a textbook attenuation in the transmission channel (Kasperson et al., 1988). Crucially, however, these misses were almost never catastrophic: the high within-one-category accuracy shows that degraded water was downgraded to Good far more often than mistaken for Excellent, so the errors clustered on adjacent categories (Excellent–Good, Good–Fair/Poor) rather than spanning the whole scale, mirroring the gradual transitions typical of water-quality data (Nasir et al., 2022). A full per-class breakdown for the best-performing model, the stacking ensemble, is given in Table 2: the classifier is most reliable on Excellent water and weakest on the minority Fair/Poor class, with the intermediate Good class in between the categories defined purely by thresholds are precisely those the index makes hardest to recover.

Table 2

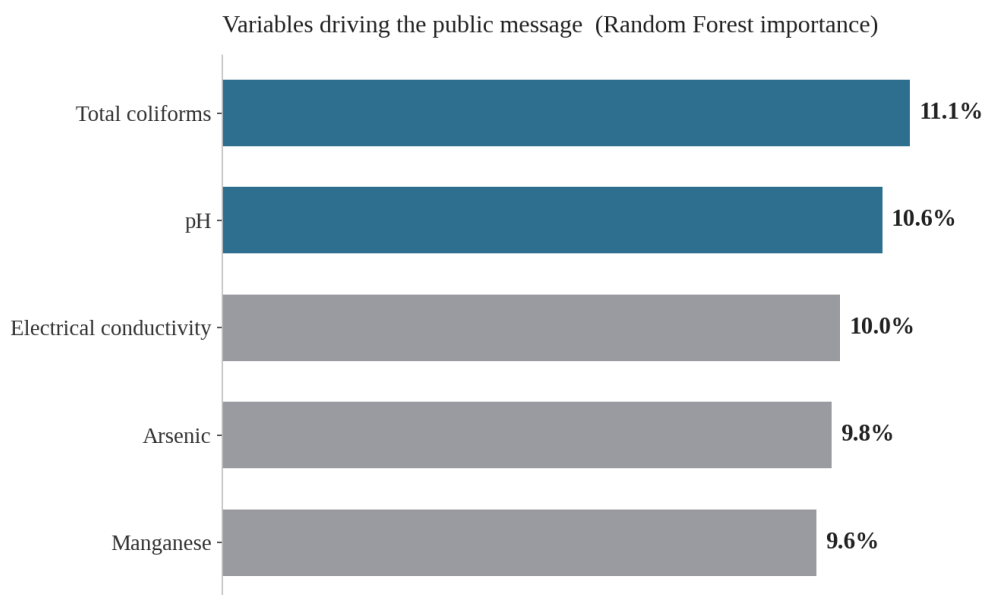
Per-class precision, recall and F1-score for the stacking ensemble (10-fold cross-validation)

Class	Precision	Recall	F1-score
Excellent	0.711	0.741	0.725
Good	0.598	0.603	0.601
Fair/Poor	0.577	0.512	0.543

Feature-importance analysis identified total coliforms, pH, electrical conductivity, arsenic and manganese as the strongest predictors, although no single variable dominated: the five leading signals each carried roughly comparable weight, around a tenth of the total (Figure 2). In framing terms, these are the variables that effectively author the public message: the category a community receives is driven by a handful of signals, several of which arsenic and manganese reflect the region's volcanic geochemistry rather than human activity (Biamont-Rojas et al., 2023; Salas-Ávila et al., 2021). The distinction matters socially, because a Fair/Poor verdict of largely geogenic origin calls for a different communicative and policy response than one driven by wastewater or mining, yet the bare category label conveys none of it.

Figure 2

Variables driving the public message: relative importance of the five strongest predictors in the Random Forest classifier



This ceiling proved robust to the choice of method. Beyond the five classifiers in Table 1, we also evaluated neural-network, instance-based and regression-to-threshold learners, alternative resampling schemes, an extensive randomized hyperparameter search and a nested cross-validation, together with features expressing each measurement as its exceedance of the national environmental-quality standards (Ministerio del Ambiente, 2017). None of these strategies surpassed roughly 65% exact accuracy, and the bias-free nested estimate was lower still (about 62%). The stability of this limit across markedly different learning approaches indicates that the barrier to recovering the exact category is informational rather than methodological: the distinction between adjacent categories above all between Good and Excellent is not encoded in the raw measurements but introduced by the index's own thresholding. Tellingly, the regression-to-threshold model, which predicts the continuous score before binning it, achieved the highest within-one-category accuracy (94.5%), underscoring that the recoverable signal is ordinal and continuous rather than sharply categorical.

What the categories leave out

The unsupervised analysis exposes the cost of legibility, and the larger dataset makes that cost unusually vivid. The reason is visible in the raw scores themselves (Figure 3): all 189 Excellent observations sit exactly at the index ceiling of 100, while 120 of the 131 Good observations pile up at a score of roughly 94.2 a hair below the Excellent cut-off of 95. Almost a third of the data therefore crowds the narrow band where one category turns into another, so that near-identical water can receive opposite public verdicts. The categories themselves bear this out: their silhouette coefficient in standardized feature space is slightly negative (-0.03): across all three categories, most observations lie as close to a neighbouring category as to their own (Figure 4), so the public categories are not natural groupings discovered in the data but cuts imposed upon it. Principal-component analysis points the same way: the first three components explain just 33.45% of total variance, and ten of the fourteen are needed to reach 80%, so no single process dominates and quality instead varies along a high-dimensional continuum (Figure 5).

Figure 3
Distribution of CCME-WQI scores

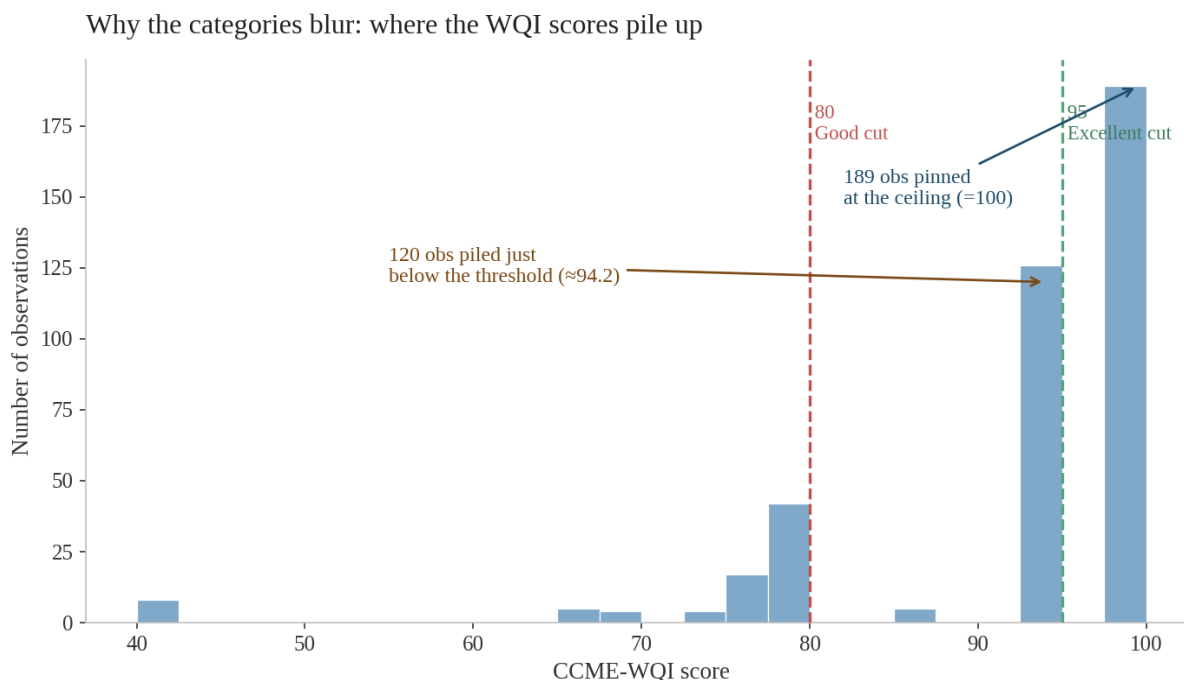


Figure 4

Silhouette plot of the three CCME-WQI categories in standardized feature space

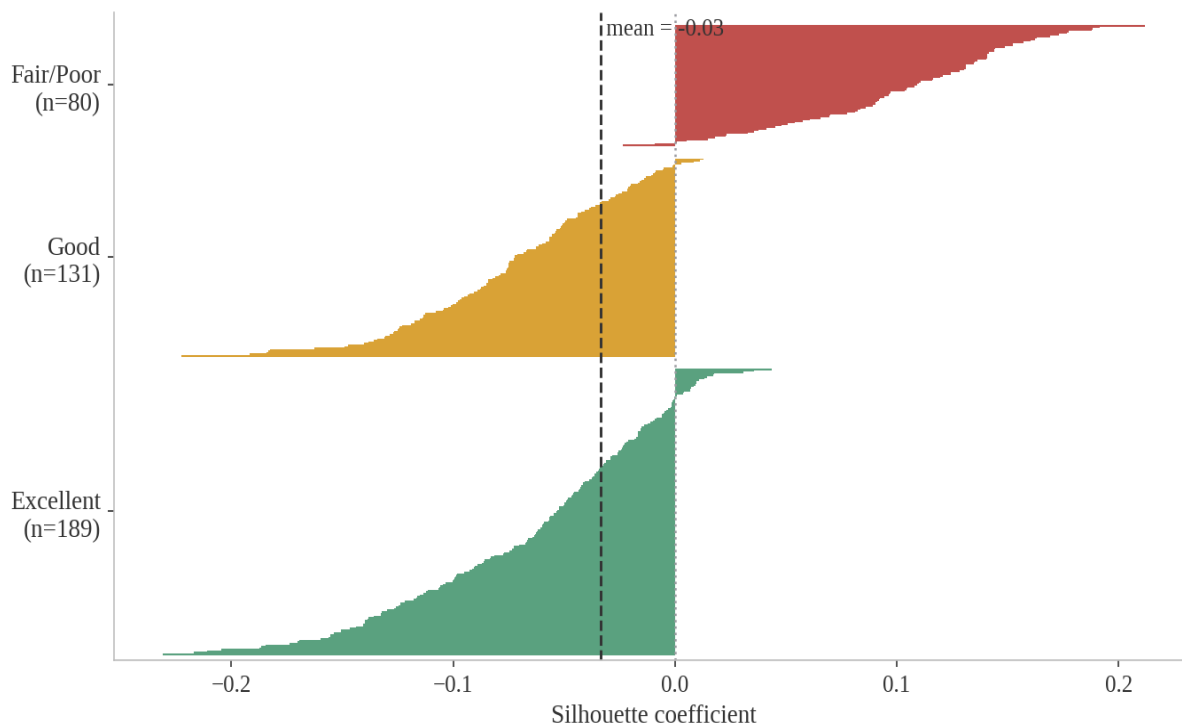
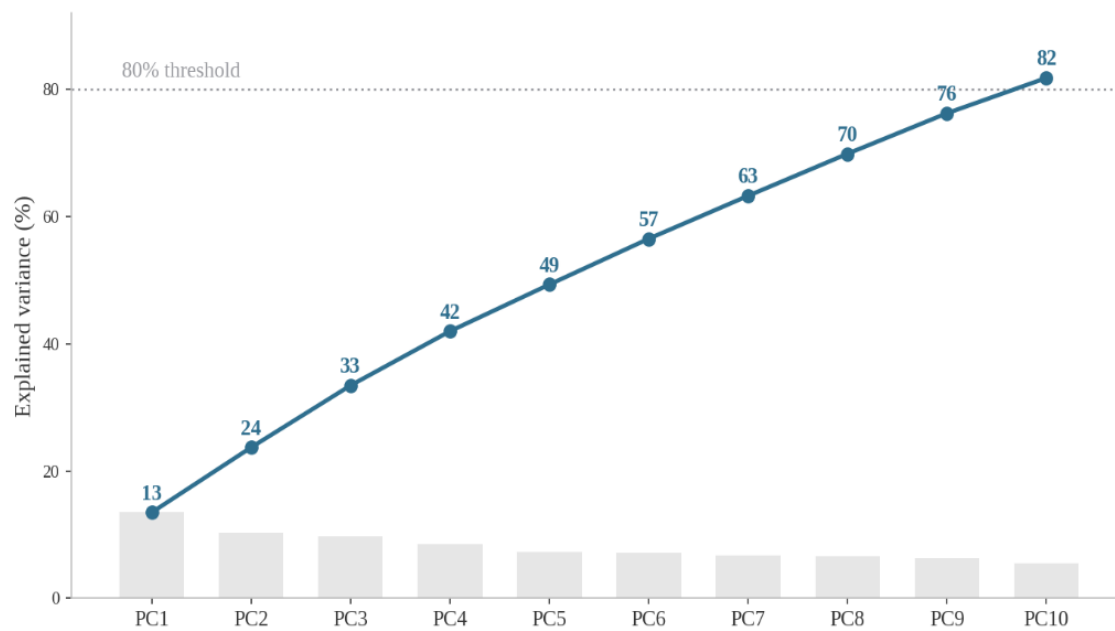


Figure 5

Variance explained by successive principal components and cumulative variance



This is the heart of the matter. The categories are meaningful but only loosely anchored in the data: the index earns its communicability by cutting sharp lines through a continuous, multidimensional reality, and the more observations there are, the more of them fall into the blurred zones around those lines. Two water bodies on opposite sides of a threshold receive opposite public messages despite near-identical conditions, while two bodies within the same category may differ substantially. The simplification that makes the message intelligible is also what makes it partial (Star & Griesemer, 1989; Fischhoff, 1995). Recognizing the index as a frame rather than a fact invites a more reflexive practice: reporting the underlying score and its uncertainty alongside the category, and flagging near-threshold observations as ambiguous rather than definitive.

Communication and society in the Altiplano

These results bear directly on how water risk is communicated in the Ilave Basin. The message must first reach its audience in an interpretable form. Monitoring data produced by the ANA are technically public yet practically inaccessible to many Quechua- and Aymara-speaking water users, for reasons of language, literacy, format and connectivity; treating environmental information as genuinely public therefore requires translation linguistic and conceptual rather than mere disclosure (Wynne, 1992; Irwin, 1995). The framing of the message also shapes its political life: whether a basin is reported as Good or as Fair/Poor conditions investment, consultation and the perceived legitimacy of community demands, so the choice of thresholds is itself a distributive decision (Entman, 1993; Massa et al., 2024).

Latin American communication theory sharpens this point. The index does not travel from the laboratory to the community as a transparent signal; it is reworked through social mediations the institutions, languages and everyday practices through which publics make technical messages their own (Martín-Barbero, 1987). In the Altiplano these mediations are hybrid: official physicochemical categories coexist with communal and traditional understandings of water and its uses, so the public message enters a cultural field that already interprets water quality on its own terms rather than a blank slate (García Canclini, 1990). Seen this way, the question is not merely whether the category is accurate but who is made visible or invisible by it — a matter of data justice, understood as fairness in how communities are represented and treated through data they did not produce and cannot easily contest (Taylor, 2017; Dencik et al., 2016). When monitoring archives are legible only to specialists, the communities most exposed to water risk are precisely those least able to question how the categories that govern their rivers are drawn.

Finally, the gap between continuous data and discrete message opens space for participatory and citizen-science approaches in which communities help interpret and contextualize indices rather than simply receive them (Irwin, 1995; Bucchi & Trench, 2014). Automated classification can support this work by rendering large monitoring archives legible and by reliably surfacing the high-risk cases that most warrant attention provided the model is chosen and tuned, as the class-rebalanced learners were here, for sensitivity to the warning category rather than for headline accuracy alone (Sandoval et al., 2025; Zhu et al., 2022). The broader lesson is that the design of an environmental index is a communication decision with social consequences, and ought to be evaluated as one.

Conclusions

Reframing automated water-quality classification as environmental-risk communication shows that water-quality indices are messages as much as measurements. Using ANA data from the Ilave River Basin, we found that the public categories of the CCME-WQI are only moderately recoverable from raw physicochemical data in exact terms: a stacking ensemble reached 65.0% accuracy, a ceiling that neither class rebalancing nor gradient boosting could break. Because the categories are ordinal, however, the same models attained about 91% within-one-category accuracy they almost never confuse excellent water with degraded so the labels track real structure even where they cannot pin the exact boundary. The category most consequential for public warning Fair/Poor nonetheless remained the hardest to detect, with even the best model recovering only about half of the worst cases. The cause is structural rather than algorithmic: nearly half the observations are pinned at the index ceiling and most of the rest crowd just below a category threshold, while a near-zero category silhouette and PCA confirm that quality varies along a continuous, high-dimensional gradient the three labels cut across rather than follow. Legibility, in short, is purchased at the cost of fidelity and the more data accumulate, the steeper that price becomes.

Three implications follow for communication and society. The framing embedded in index thresholds shapes how publics perceive water risk and so deserves deliberate, reflexive design; environmental information becomes genuinely public only when it is translated into accessible, multilingual forms for the communities that depend on it;

and data-driven classification, used judiciously, can make monitoring legible and amplify rather than suppress the signals that matter most. Future work should pair this analysis with audience-reception studies in Quechua- and Aymara-speaking communities, test multilingual and visual formats for communicating water risk, and explore participatory indicators co-designed with water users, so that the passage from technical data to public message broadens rather than narrows the public conversation about water in the Altiplano.

References

- Abbas, F., Cai, Z., Shoaib, M., Iqbal, J., Ismail, M., Arifullah, Alrefaei, A. F., & Albeshr, M. F. (2024). Machine learning models for water quality prediction: A comprehensive analysis and uncertainty assessment in Mirpurkhas, Sindh, Pakistan. *Water*, 16(7), 941. <https://doi.org/10.3390/w16070941>
- Autoridad Nacional del Agua. (2016). Protocolo nacional para el monitoreo de la calidad de los recursos hídricos superficiales (Resolución Jefatural N°010-2016-ANA). Autoridad Nacional del Agua.
- Banda, T. D., & Kumarasamy, M. V. (2020). Development of water quality indices (WQIs): A review. *Polish Journal of Environmental Studies*, 29(3), 2011–2021. <https://doi.org/10.15244/pjoes/110526>
- Beck, U. (1992). *Risk society: Towards a new modernity*. SAGE Publications.
- Bentéjac, C., Csörgő, A., & Martínez-Muñoz, G. (2021). A comparative analysis of gradient boosting algorithms. *Artificial Intelligence Review*, 54(3), 1937–1967. <https://doi.org/10.1007/s10462-020-09896-5>
- Biamont-Rojas, I. E., Cardoso-Silva, S., Figueira, R. C. L., Kim, B. S. M., Alfaro-Tapia, R., & Pompêo, M. (2023). Spatial distribution of arsenic and metals suggest a high ecotoxicological potential in Puno Bay, Lake Titicaca, Peru. *Science of the Total Environment*, 871, 162051. <https://doi.org/10.1016/j.scitotenv.2023.162051>
- Biau, G., & Scornet, E. (2016). A random forest guided tour. *TEST*, 25(2), 197–227. <https://doi.org/10.1007/s11749-016-0481-7>
- Bucchi, M., & Trench, B. (Eds.). (2014). *Routledge handbook of public communication of science and technology* (2nd ed.). Routledge. <https://doi.org/10.4324/9780203483794>
- Cervantes, J., Garcia-Lamont, F., Rodríguez-Mazahua, L., & Lopez, A. (2020). A comprehensive survey on support vector machine classification: Applications, challenges and trends. *Neurocomputing*, 408, 189–215. <https://doi.org/10.1016/j.neucom.2019.10.118>
- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic minority over-sampling technique. *Journal of Artificial Intelligence Research*, 16, 321–357. <https://doi.org/10.1613/jair.953>
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794). <https://doi.org/10.1145/2939672.2939785>
- Chen, W., Xu, D., Pan, B., Zhao, Y., & Song, Y. (2024). Machine learning-based water quality classification assessment. *Water*, 16(20), 2951. <https://doi.org/10.3390/w16202951>
- Dencik, L., Hintz, A., & Cable, J. (2016). Towards data justice? The ambiguity of anti-surveillance resistance in political activism. *Big Data & Society*, 3(2), 1–12. <https://doi.org/10.1177/2053951716679678>
- Entman, R. M. (1993). Framing: Toward clarification of a fractured paradigm. *Journal of Communication*, 43(4), 51–58. <https://doi.org/10.1111/j.1460-2466.1993.tb01304.x>
- Escobar-Mamani, F., Moreno-Terrazas, E., Siguyro-Mamani, H., & Argota-Pérez, G. (2023). Physicochemical characterization and presence of heavy metals in the trout farming area of Lake Titicaca, Peru. *Sains Tanah – Journal of Soil Science and Agroclimatology*, 20(2), 140–149. <https://doi.org/10.20961/stjssa.v20i2.62357>
- Fischhoff, B. (1995). Risk perception and communication unplugged: Twenty years of process. *Risk Analysis*, 15(2), 137–145. <https://doi.org/10.1111/j.1539-6924.1995.tb00308.x>
- Frincu, R. M. (2025). Artificial intelligence in water quality monitoring: A review of water quality assessment applications. *Water Quality Research Journal*, 60(1), 164–176. <https://doi.org/10.2166/wqrj.2024.049>
- García Canclini, N. (1990). *Culturas híbridas: Estrategias para entrar y salir de la modernidad*. Grijalbo.
- Irwin, A. (1995). *Citizen science: A study of people, expertise and sustainable development*. Routledge.
- Kasperson, R. E., Renn, O., Slovic, P., Brown, H. S., Emel, J., Goble, R., Kasperson, J. X., & Ratick, S. (1988). The social amplification of risk: A conceptual framework. *Risk Analysis*, 8(2), 177–187. <https://doi.org/10.1111/j.1539-6924.1988.tb01168.x>

- Krishnaraj, A., & Deka, P. C. (2020). Spatial and temporal variations in river water quality of the Middle Ganga Basin using unsupervised machine learning techniques. *Environmental Monitoring and Assessment*, 192(12), 744. <https://doi.org/10.1007/s10661-020-08624-4>
- Martín-Barbero, J. (1987). *De los medios a las mediaciones: Comunicación, cultura y hegemonía*. Gustavo Gili.
- Massa, A., & Comunello, F. (2024). Natural and environmental risk communication: A scoping review of campaign experiences, applications and tools. *International Journal of Disaster Risk Reduction*, 114, 104936. <https://doi.org/10.1016/j.ijdrr.2024.104936>
- Ministerio del Ambiente. (2017). *Estándares de calidad ambiental (ECA) para agua* (D.S. N° 004-2017-MINAM). Ministerio del Ambiente. <https://www.gob.pe/institucion/minam/normas-legales/3671-004-2017-minam>
- Mogane, L. K., Masebe, T., Msagati, T. A. M., & Ncube, E. (2023). A comprehensive review of water quality indices for lotic and lentic ecosystems. *Environmental Monitoring and Assessment*, 195(8), 926. <https://doi.org/10.1007/s10661-023-11512-2>
- Nasir, N., Kansal, A., Alshaltone, O., Barneih, F., Sameer, M., Shanableh, A., & Al-Shamma'a, A. (2022). Water quality classification using machine learning algorithms. *Journal of Water Process Engineering*, 48, 102920. <https://doi.org/10.1016/j.jwpe.2022.102920>
- Renn, O. (2008). *Risk governance: Coping with uncertainty in a complex world*. Earthscan.
- Salas-Ávila, D., Chaíña-Chura, F. F., Belizario-Quispe, G., Quispe-Mamani, E., Huanqui-Pérez, R., Velarde-Coaquira, E., Bernedo-Colca, F., Salas-Mercado, D., & Hermoza-Gutiérrez, M. (2021). Evaluación de metales pesados y comportamiento social asociados con la calidad del agua en el río Suches, Puno, Perú. *Tecnología y Ciencias del Agua*, 12(6), 145–195. <https://doi.org/10.24850/j-tyca-2021-06-04>
- Sandoval, S., Bui, J., & Hopfer, S. (2025). Wildfire and smoke risk communication: A systematic literature review from a health equity focus. *International Journal of Environmental Research and Public Health*, 22(3), 368. <https://doi.org/10.3390/ijerph22030368>
- Scikit-learn developers. (2024). *Scikit-learn: Machine learning in Python (Version 1.5)*. <https://scikit-learn.org/stable/modules/clustering.html>
- Star, S. L., & Griesemer, J. R. (1989). Institutional ecology, 'translations' and boundary objects: Amateurs and professionals in Berkeley's Museum of Vertebrate Zoology, 1907–1939. *Social Studies of Science*, 19(3), 387–420. <https://doi.org/10.1177/030631289019003001>
- Swain, L. G. (2024). The evolution of the Canadian (CCME) Water Quality Index. *Water Quality Research Journal*, 59(4), 342–346. <https://doi.org/10.2166/wqrj.2024.052>
- Taylor, L. (2017). What is data justice? The case for connecting digital rights and freedoms globally. *Big Data & Society*, 4(2), 1–14. <https://doi.org/10.1177/2053951717736335>
- Wolpert, D. H. (1992). Stacked generalization. *Neural Networks*, 5(2), 241–259. [https://doi.org/10.1016/S0893-6080\(05\)80023-1](https://doi.org/10.1016/S0893-6080(05)80023-1)
- Wynne, B. (1992). Misunderstood misunderstanding: Social identities and public uptake of science. *Public Understanding of Science*, 1(3), 281–304. <https://doi.org/10.1088/0963-6625/1/3/004>
- Zhu, M., Wang, J., Yang, X., Zhang, Y., Zhang, L., Ren, H., Wu, B., & Ye, L. (2022). A review of the application of machine learning in water quality evaluation. *Eco-Environment & Health*, 1(2), 107–116. <https://doi.org/10.1016/j.eehl.2022.06.001>

AUTHOR CONTRIBUTIONS (CREDIT)

The authors confirm their contributions to the paper as follows. Pedro Alvaro Edwin Gallegos Pasco: conceptualization, methodology, formal analysis, writing – original draft, and project administration. Haydee Celia Pineda Chaíña: conceptualization, supervision, and writing – review and editing. Luz Marina Teves Ponce: investigation, data curation, and writing – review and editing. Nestor Omar Mercado Ayamamani: methodology, software, and validation. Alfredo Mamani Canqui: investigation, resources, and data curation. Johan Denis Aguilar Ramirez: software, visualization, and writing – review and editing. All authors reviewed the results and approved the final version of the manuscript.

CONFLICTS OF INTEREST

The authors declare no conflict of interest.