

PREVISÃO SUPERVISIONADA DE VOLUME COMO SUBSÍDIO À IDENTIFICAÇÃO DE SECAS NO SISTEMA CANTAREIRA, SÃO PAULO

Raphael Ferreira Perez¹ ; Veronica Lima Gonsalez² & Darwin Hermilo Cora Quispe³ & Joaquin Ignacio Bonnacarrère⁴

RESUMO – A previsão de condições futuras de armazenamento em reservatórios é uma ferramenta relevante para antecipar situações de escassez hídrica e apoiar o planejamento de sistemas de abastecimento. Este trabalho avalia o uso de aprendizado de máquina supervisionado para estimar o volume armazenado no Sistema Cantareira, com foco inicial no reservatório Jaguari/Jacareí, principal componente de armazenamento do sistema. Foram utilizadas séries mensais de precipitação, evapotranspiração, *total water storage*, vazão natural, volume armazenado e variáveis operacionais entre jan/1984 e mar/2024. A partir dessas séries, foram construídos atributos temporais por meio de técnicas de *feature engineering*, incluindo defasagens, médias móveis, variáveis sazonais e índices padronizados de seca. Três conjuntos de atributos foram avaliados com o algoritmo XGBoost para previsão do volume em horizonte de três meses. O melhor experimento apresentou NSE de 0,675 e correlação de 0,859 no período de teste do modelo. A análise de interpretabilidade SHAP (*SHapley Additive exPlanations*) indicou maior influência dos atributos: volume atual, vazões operacionais para abastecimento e vazão natural. Os resultados preliminares sugerem que o modelo proposto é promissor como ferramenta de apoio à gestão de recursos hídricos.

Palavras-Chave – previsão de volume, XGBoost, seca hidrológica, escassez hídrica

ABSTRACT – Forecasting future reservoir storage conditions is relevant for anticipating water scarcity and supporting water supply planning. This study evaluates the use of supervised machine learning to estimate stored volume in the Cantareira System, focusing initially on the Jaguari/Jacareí reservoir, the main storage component of the system. Monthly time series of precipitation, evapotranspiration, total water storage, natural inflow, stored volume, and operational variables from Jan/1984 to Mar/2024 were used. Temporal predictors were derived through feature engineering techniques, including lagged variables, moving-window statistics, seasonal attributes, and standardized drought indices. Three predictor sets were tested using the XGBoost algorithm to forecast reservoir storage at a three-month horizon. The best experiment achieved an NSE of 0.675 and a correlation of 0.859 in the test period. SHAP (*SHapley Additive exPlanations*) interpretability analysis indicated greater influence from predictors: current storage, operational outflows for water supply, and natural outflows from rainfall-runoff. Preliminary results suggest that the developed model is promising as a tool for water resources management.

Key-words – storage forecast, XGBoost, hydrological drought, water scarcity

1) Afiliação: Fundação Centro Tecnológico de Hidráulica e Departamento de Engenharia Hidráulica e Ambiental, Escola Politécnica da Universidade de São Paulo, São Paulo (SP), 05508-010, Brasil, rfperez@usp.br

2) Afiliação: Departamento de Engenharia Hidráulica e Ambiental, Escola Politécnica da Universidade de São Paulo, São Paulo (SP), 05508-010, Brasil, veronica.gonsalez@usp.br

3) Afiliação: Departamento de Engenharia Hidráulica e Ambiental, Escola Politécnica da Universidade de São Paulo, São Paulo (SP), 05508-010, Brasil, darwin.cora@usp.br

4) Afiliação: Departamento de Engenharia Hidráulica e Ambiental, Escola Politécnica da Universidade de São Paulo, São Paulo (SP), 05508-010, Brasil, joaquinbonne@usp.br

INTRODUÇÃO

A escassez hídrica representa uma condição em que os recursos hídricos disponíveis são insuficientes para atender às demandas médias de longo prazo da sociedade e dos ecossistemas (EU, 2007). Embora fatores naturais, como a variabilidade climática e a redução das precipitações, estejam diretamente associados à ocorrência de escassez hídrica, fatores antrópicos como crescimento populacional, urbanização acelerada, aumento das demandas por água e deficiências na gestão dos recursos hídricos intensificam significativamente esse problema.

Entre os principais fatores associados à escassez hídrica está a ocorrência de períodos prolongados de precipitação abaixo da média histórica. A crise hídrica ocorrida entre 2014 e 2015 no Sudeste do Brasil evidenciou a vulnerabilidade dos sistemas de abastecimento em relação à precipitação. Na Região Metropolitana de São Paulo, precipitações abaixo da média e temperaturas elevadas levaram reservatórios estratégicos, como o Sistema Cantareira, a níveis críticos de armazenamento (NOBRE *et al.*, 2016). Além dos impactos no abastecimento urbano, o evento gerou prejuízos econômicos e aumentou a pressão sobre serviços essenciais (MUNICH RE, 2015), reforçando a necessidade de ferramentas capazes de antecipar cenários de escassez hídrica e apoiar a gestão dos recursos hídricos.

Nesse contexto, a previsão de secas e da disponibilidade hídrica permanece como um desafio científico e operacional. Os modelos utilizados para previsão de secas podem ser agrupados em três categorias principais: modelos dinâmicos baseados em processos físicos, modelos estatísticos orientados por dados e modelos híbridos que combinam ambas as abordagens (AGHAKOUCHAK *et al.*, 2022). Embora modelos físicos sejam amplamente empregados, sua aplicação frequentemente depende de grandes volumes de dados e elevada capacidade computacional, o que limita sua utilização em regiões com baixa disponibilidade de informações hidrometeorológicas (AHMED *et al.*, 2024).

As abordagens baseadas em aprendizado de máquina têm recebido crescente atenção na área de recursos hídricos. Esses modelos apresentam elevada flexibilidade para integrar diferentes fontes de dados e capacidade de identificar relações não lineares complexas entre variáveis hidrológicas, climáticas e ambientais. Aplicações recentes incluem previsão de vazões superficiais, monitoramento de aquíferos, previsão de índices de seca e suporte à gestão de águas subterrâneas (AKBARIAN; SAGHAFIAN; GOLIAN, 2023; CHENG *et al.*, 2020; KOCHHAR *et al.*, 2022; TEIMOORI; OLYA; MILLER, 2023). No Brasil, estudos como os de Galleir *et al.* (2025) e Santos *et al.* (2009)

demonstraram resultados promissores na aplicação de técnicas de machine learning para previsão de secas agrícolas e meteorológicas em diferentes regiões do país.

Apesar dos avanços recentes, ainda existem desafios relacionados à construção de modelos robustos apropriados para o complexo sistema de abastecimento público com armazenamento e transposições. Nesse sentido, torna-se fundamental integrar variáveis do meio físico, demandas hídricas e captações de água em modelos preditivos orientados por impacto. A integração dessas informações pode fornecer suporte técnico para ações de mitigação, planejamento e adaptação em sistemas de abastecimento na Região Metropolitana de São Paulo (RMSP).

O objetivo deste trabalho abrange a utilização de técnicas de aprendizado de máquina, utilizando o algoritmo XGBoost, para estimar estados de volume de curto período e assim identificar tendências de escassez hídrica. Como objeto de estudo, utilizou-se o reservatório Jaguari-Jacareí no Sistema Cantareira, principal componente de sistema de abastecimento público paulistano e essencial para o desenvolvimento socioeconômico da RMSP.

METODOLOGIA

O Sistema Cantareira (SC), composto por cinco reservatórios interligados (Jaguari, Jacareí, Cachoeira, Atibainha e Paiva Castro), possui um volume útil total de aproximados 981 bilhões de litros e atende 46% de toda a população que compõe a Região Metropolitana de São Paulo (RMSP) (ANA, 2026). Como ilustra a Figura 1 (com a qual nomeamos Jaguari-Jacareí de “JAGJAC”, Cachoeira de “CACH”, Atibainha de “ATIB” e Paiva Castro de “PCASTR”), a comunicação entre os reservatórios garante distribuição de água à cidade de São Paulo e a municípios vizinhos. A relevância socioeconômica da região, aliada a um histórico recente de escassez hídrica, impulsiona a investigação de técnicas que aprimorem a gestão e o processo decisório de planejamento.

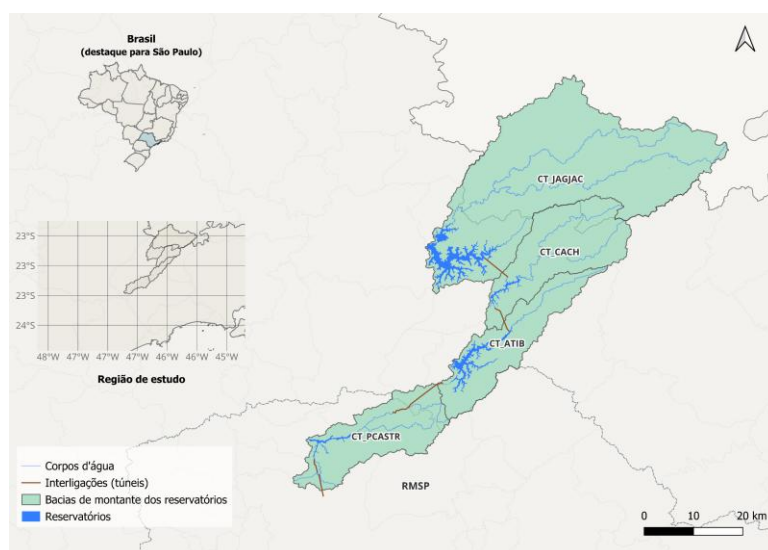


Figura 1 – Localização do Sistema Cantareira e seus reservatórios comunicantes

Atualmente a gestão de recursos hídricos é ancorada no monitoramento e na modelagem hidrológica previstas no Sistema de Suporte à Decisão (SSD) da Sabesp. Dados do meio físico estão organizados na plataforma, que também dispõe de modelos de chuva-vazão para prever vazões futuras a partir de previsões meteorológicas. Para transformar dados de vazão em volume armazenado, é necessário um modelo de alocação que otimize a operação com todas as restrições físicas e regulatórias da infraestrutura atualizadas. A complexidade da rede e as incertezas inerentes aos modelos inviabiliza a previsão de volumes armazenados com as ferramentas atuais do SSD Sabesp.

Para superar este desafio, o uso de aprendizado de máquina em recursos hídricos, ou *machine learning*, propõe determinar métodos que complementem a prática de hidrologia tradicional com ferramentas computacionais e estatísticas para extrair informações diretamente de conjuntos de dados (GHOBADI; KANG, 2023). Modelos de *machine learning* (ML) são capazes de identificar relações não lineares entre entradas e saídas hidrológicas para produzir previsões confiáveis, mesmo não conhecendo o comportamento físico explícito que governa tais variáveis (AHMED *et al.*, 2024).

Em meio às abordagens de ML, algoritmos supervisionados abrangem modelos para previsão, ou estimativa, de uma saída a partir de uma ou mais entradas (JAMES *et al.*, 2023). Em termos hidrológicos, conta-se com um conjunto de dados hidroclimáticos de *input* (precipitação, evapotranspiração ou vazão, por exemplo) e uma variável específica de interesse para *output* (volume armazenado), em uma série temporal suficientemente grande para seccionar o experimento de regressão em treinamento (que corresponde a etapa de calibração) e teste (validação) do modelo.

As árvores de decisão são mecanismos supervisionados para regressão em hidrologia, que realizam previsões a partir de sucessivas segmentações do espaço definido pelos *inputs* (BREIMAN, 1984; MORETTIN; SINGER, 2025). A eficiência da previsão é aprimorada, e a variância de respostas é contida por meio de processos mais sofisticados, como o algoritmo *gradient boosting*, que adota múltiplas árvores menores iterativamente (FRIEDMAN, 2001). Especificamente em modelagem de secas, nota-se uma crescente de estudos favoráveis com metodologias que aplicam o algoritmo XGBoost (ZHANG *et al.*, 2023; ALI *et al.*, 2023), um desdobramento de *gradient boosting* escolhido para suportar as previsões para o Sistema Cantareira neste estudo.

A fim de construir uma metodologia robusta e de fácil interpretação, o *output* foi definido como o volume futuro do reservatório, em um horizonte de 3 meses. As projeções volumétricas via ML permitem que tomadores de decisão orientem para ações de mitigação baseadas na regulação do Protocolo de Secas (SP Águas, 2025).

A Figura 2 apresenta os grupos de *features* de previsão consideradas inicialmente na construção do modelo. Tomando séries mensais de precipitação, evapotranspiração, *total water storage* (representação de dados de satélite para volumes de água sub e superficial terrestre), vazão natural e volume para a região de estudo, variando de jan/1984 a mar/2024, listou-se eventuais transformações para construir vetores de sazonalidade e memória temporal (como defasagens, médias móveis e índices padronizados consagrados: SPI, SPEI e SRI, por exemplo), necessários para o aprendizado de máquina e a previsão supervisionada a partir de séries históricas. Este pré-processamento é comumente chamado de *feature engineering* (KUHN; JOHNSON, 2019).

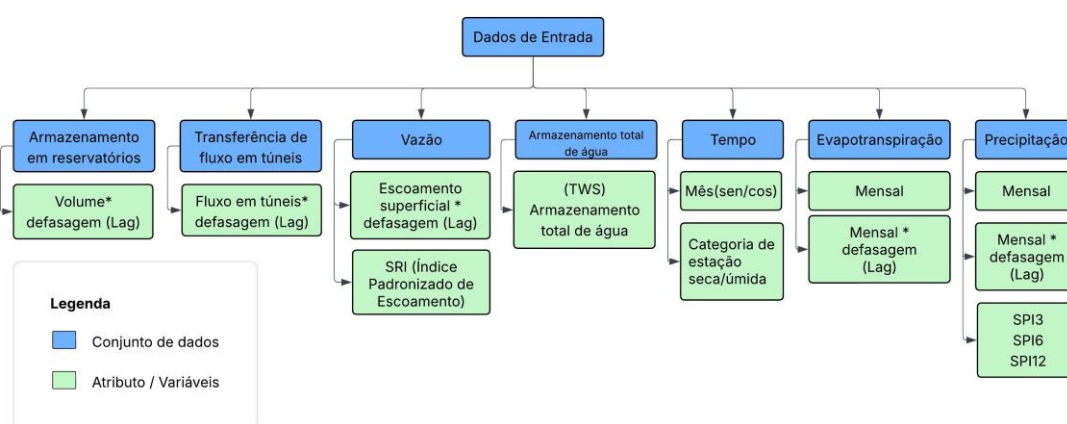


Figura 2 – Conjuntos de *features* exploradas como preditoras de modelagem

A seleção dos atributos finais foi baseada na análise da correlação de Pearson entre as variáveis preditoras e a variável objetivo (volume armazenado nos reservatórios) em duas etapas: exclusão de variáveis altamente correlacionadas e eliminação baseada na importância das variáveis. Esses procedimentos foram adotados com o objetivo de simplificar o modelo de previsão, além de reduzir o tempo e o esforço computacional necessários nas etapas posteriores de treinamento e validação do algoritmo (FERREIRA *et al.*, 2021). A partir dos valores de correlação, são eliminadas variáveis com correlação de Pearson superior a 90% com índice de multicolinearidade, e selecionadas as variáveis com maiores correlações com a variável resposta, que é o volume armazenado no mês.

A inclusão das features ao modelo foi realizada de forma parcimoniosa, comparando em 3 etapas se houve avanço na *performance* do modelo a partir da inclusão de um grupo de features. Com isto, pôde-se comparar melhorias de previsão com a) diferentes ganhos de correlação com observações de teste e b) diminuição de erros médios quadráticos. A inclusão de features em etapas foi observada a partir de 3 experimentos:

- Experimento A: dados de chuva, evapotranspiração, vazão natural, volume, “lags menos 1” destas quatro (ou seja, dados de véspera) e *inputs* de definição temporal (senos e cossenos para representação contínua de meses e rótulo de período úmido e seco para cada mês)

- Experimento B: *features* de A mais vazões operacionais pertinentes ao reservatório: descarga a jusante, vazão veiculada no túnel 7 em direção ao Cachoeira e “lags menos 1” de ambas
- Experimento C: *features* de A mais inputs selecionados por correlação: média móvel de 6 meses de vazão, TWS, SRI₁₂ (*Standard Runoff Index* com janela de 12 meses) e média acumulada de evapotranspiração

RESULTADOS E DISCUSSÃO

Para o reservatório Jaguari/Jacareí, o modelo apresentou elevada correlação entre volume previsto e observado, baixo erro absoluto (MAE) e quadrático médio (RMSE) e coeficiente de Nash-Sutcliffe (NSE) próximo de 1, resumidos na Tabela 1. Os experimentos foram comparados entre si e com a Persistência. No contexto de ML para séries temporais, persistência é uma abordagem de referência em que o valor futuro é assumido como igual ao valor observado mais recente, servindo como baseline para avaliar se o modelo realmente possui capacidade preditiva.

Tabela 1 – Métricas de performance para experimentos em Jaguari/Jacareí para partição de teste (jan/2019 a mar/2024)

Experimento	RMSE	MAE	Correlação	NSE
Persistência	120.50			0.377
B	87.04	68.65	0.859	0.675
A	87.23	64.98	0.883	0.673
C	100.19	78.20	0.852	0.569

Fonte: Autores

Apesar de haver um *trade-off* em termos de correlação (confrontando 0.859 frente a 0.883), a Tabela 1 indica que o experimento com melhor desempenho preditivo é B, com um *mix* de variáveis temporais, climáticas, hidrológicas e operacionais. A seleção de B condiz com a própria natureza intrínseca ao estudo de caso, uma vez que o Sistema Cantareira (aqui representado por seu componente de maior dimensão, a Represa Jaguari/Jacareí) possui complexidade tal que um modelo de aprendizado de máquina não pode ser unicamente alimentado com *features* hidroclimáticas. Para que as previsões sejam significativas, conforme métricas de erro médio quadrático, é imperativo descrever também a conjuntura hidráulica e operacional que governa o Sistema, via balanço por meio de descargas e transferências internas (túnel 7).

As previsões de volume de horizonte de 3 meses aferidas por ML foram comparadas com os volumes observados na partição de teste dos dados (Figura 3). De forma geral houve concordância entre volume simulado e observado, com maior dificuldade do modelo de representar volumes elevados acima de 600 hm³.

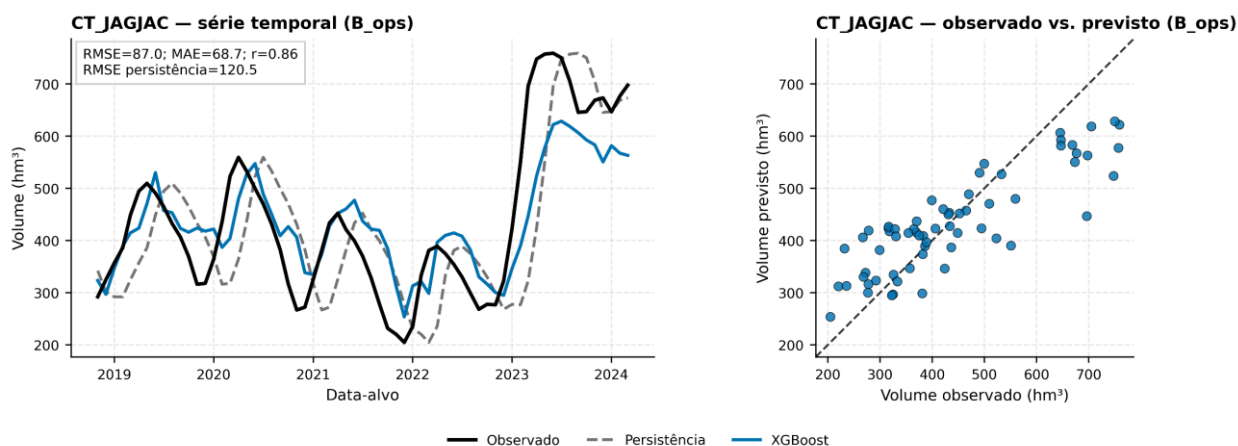


Figura 3 – Verificação de performance do modelo no *split* de teste para melhor experimento

A Figura 3 compara três séries de dados no intervalo jan/2019 – mar/2024 (visto que o período 1984 – 2018 foi reservado para treinamento e validação). Mesmo com uma leve defasagem de aproximadamente 1 mês (via descolamento horizontal), nota-se que as previsões possuem elevada correlação (aproximados 86%) e comportamento linear quando comparadas com os dados observados (conforme gráfico de dispersão). O erro RMSE atribuído ao modelo é cerca de 33 hm³ inferior à persistência, com um indicador NSE de 0.675 no teste, o que ratifica a hipótese de que o algoritmo de aprendizado de máquina é capaz de descrever volume futuro e com esta estimativa indicar estados de seca de curto prazo.

É pertinente destacar que rotinas tradicionais de planejamento e previsão inevitavelmente apoiam-se em múltiplos modelos acoplados (chuva-vazão – para aferição de séries de escoamento futuro – e alocação de água – para distribuição de vazão e evolução temporal de reservas de volume frente a demandas consuntivas). Embora essas abordagens representem de forma consistente a interação entre os processos hidrológicos e o uso antrópico da água, elas envolvem elevado número de parâmetros, etapas de calibração e propagação de incertezas. Nesse contexto, abordagens baseadas em aprendizado de máquina podem representar relações complexas diretamente a partir dos dados, reduzindo a necessidade de encadeamento de modelos e o esforço computacional, sem perder a capacidade de apoiar o processo de tomada de decisão e interpretação dos resultados.

Para dissociar a utilização de algoritmos de ML de modelos *black-box*, cujas entradas e saídas não são plenamente concatenadas e explicadas, a utilização de um complemento de xAI (*eXplainable Artificial Intelligence*) como o SHAP (*SHapley Additive exPlanations*) (Figura 4) sugere quais os principais *drivers* de uma previsão, o que traz maior racionalidade à modelagem. Podemos observar que as *features* mais relevantes e influentes para a estimativa de volumes em 3 meses são: o próprio volume em *t* (ou seja, o comportamento atual dita o estado futuro, seja em termos de aumento ou

deplecionamento de reservação), a vazão transferida internamente para jusante (túnel 7 em direção ao Reservatório Cachoeira) e as variáveis que comunicam a entrada de água na bacia (vazão natural e precipitação). De modo intuitivo pode-se justificar como a previsão supervisionada é calculada, o que confere robustez e confiança em esferas de gerenciamento de recursos hídricos.

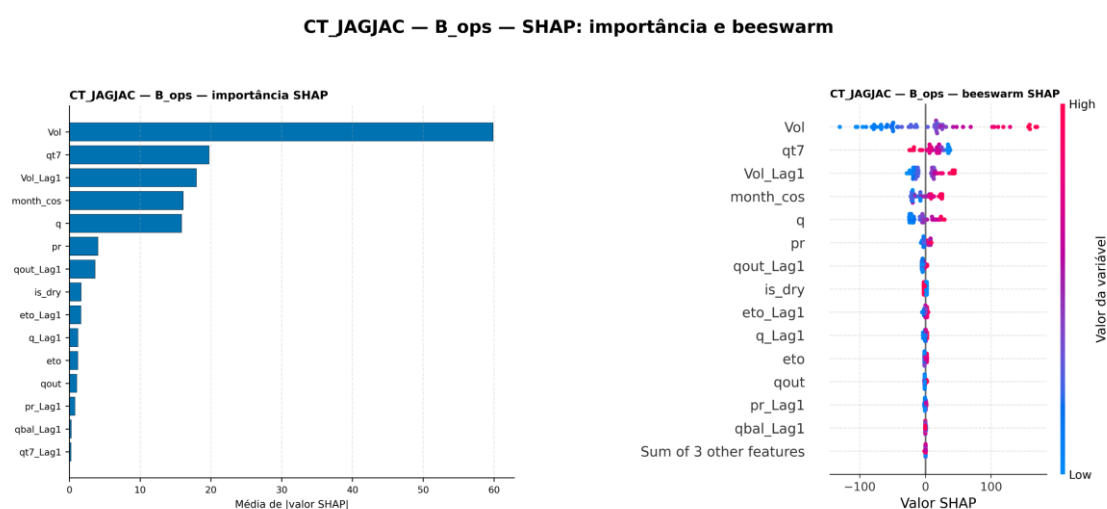


Figura 4 – Métricas SHAP ajustadas para o experimento B de CT_JAGJAC

Como perspectiva de continuidade, destaca-se que a metodologia é reproduzível em reservatórios de jusante no próprio Sistema Cantareira ou em conjunturas semelhantes, mediante disponibilidade de dados e ajuste temporal adequado para refletir a inércia e detenção características de cada reservatório. Uma possibilidade enriquecedora é modelar o SC como um sistema equivalente, por meio de uma representação que tome suas represas e as transforme em um reservatório resultante. Apesar de haver uma perda inevitável de discretização, este artifício é capaz de melhor representar o Sistema como um todo e assim abarcar todas suas fontes produtoras para abastecimento.

CONCLUSÃO

O modelo de previsão de volume armazenado tem como objetivo auxiliar na operação e gerenciamento dos recursos hídricos no reservatório Jaguari/Jacareí do Sistema Cantareira, na RMSP. Utilizando somente dados históricos, o algoritmo de aprendizado de máquina identifica padrões nos dados hidroclimáticos e operacionais para prever o volume armazenado do reservatório futuro com tempo de latência de 3 meses. A seleção dos atributos preditores foi baseada em correlações de Pearson para garantir boa acurácia de forma parcimoniosa, resultando em 17 atributos para o algoritmo de XGBoost (dado a partir de *feature engineering* das variáveis brutas de precipitação, evapotranspiração, vazão e volume).

O volume armazenado é a métrica operacional mais utilizada para aferir a capacidade de abastecimento do Sistema Cantareira, frequentemente utilizada em reportagens na mídia em porcentagem. Para obter uma previsão deste parâmetro atualmente é necessário utilizar um modelo hidrológico de chuva-vazão baseado em previsões de precipitação seguido de um modelo complexo de alocação de água que possa otimizar e simular as transferências de água no sistema de abastecimento. Este projeto propõe um novo modelo que substitui essas etapas e prevê o volume armazenado com uma metodologia de aprendizado de máquina.

O modelo do reservatório Jaguari/Jacareí apresentou previsões coerentes com os valores observados de volume armazenado, com RMSE de 87 hm³ e correlação de 86%. Estes resultados confirmam que a abordagem de aprendizado de máquina com XGBoost tem o potencial de auxiliar o gerenciamento de recursos hídricos com previsões de três meses, antecipando alertas de risco no abastecimento do maior reservatório de água potável de São Paulo.

Apesar de se tratar de um modelo simples, a abordagem é eficiente e pode se tornar mais robusta em futuros estudos com a interpretação das contribuições de cada atributo através de técnicas de inteligência artificial explicável. Os valores SHAP são uma técnica utilizada para interpretar modelos de aprendizado de máquina, quantificando a contribuição individual de cada variável preditora nas previsões do modelo. Essa abordagem permite identificar a importância e a influência de cada variável, inclusive em relações não lineares e interações complexas. Em estudos hidrológicos e ambientais, a aplicação de SHAP é relevante por aumentar a interpretabilidade e a transparência dos modelos, auxiliando na identificação das variáveis mais importantes e na avaliação da coerência física dos padrões aprendidos pelo algoritmo.

REFERÊNCIAS

- AGHAKOUCHAK, A. *et al.* Status and prospects for drought forecasting: opportunities in artificial intelligence and hybrid physical–statistical forecasting. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, v. 380, n. 2238, p. 20210288, 24 out. 2022.
- AHMED, Ashraf A. *et al.* Applications of machine learning to water resources management: A review of present status and future opportunities. *Journal of Cleaner Production*, v. 441, p. 140715, 15 fev. 2024.
- AKBARIAN, Mohammad; SAGHAFIAN, Bahram; GOLIAN, Saeed. Monthly streamflow forecasting by machine learning methods using dynamic weather prediction model outputs over Iran. *Journal of Hydrology*, v. 620, p. 129480, 1 mai. 2023.
- ALI, S. *et al.* Spatial Downscaling of GRACE Data Based on XGBoost Model for Improved Understanding of Hydrological Droughts in the Indus Basin Irrigation System (IBIS). *Remote Sensing*, vol. 15, p. 873, 3 fev. 2023.
- ANA. Sistema Cantareira. Disponível em <https://www.gov.br/ana/pt-br/sala-de-situacao/sistema-cantareira/sistema-cantareira-saiba-mais>. 13 mai. 2026.

_____. Sistema Cantareira continuará a operar na Faixa de Atenção em maio. Disponível em <https://www.gov.br/ana/>

pt-br/assuntos/noticias-e-eventos/noticias/sistema-cantareira-continuara-a-operar-na-faixa-de-atencao-em-maio?utm_source=chatgpt.com. 13 mai. 2026.

BREIMAN, L.; FRIEDMAN, J. H.; OLSHEN, R. A.; STONE, C. J. *Classification and Regression Trees*. Belmont: Wadsworth, 1984.

CHENG, M. *et al.* Long lead-time daily and monthly streamflow forecasting using machine learning methods. *Journal of Hydrology*, v. 590, p. 125376, 1 nov. 2020.

EU. *Addressing the challenge of water scarcity and droughts in the European Union: Note for the European Committee of the Regions*. [S.l.: s.n.], 2007.

FERREIRA, Renan Gon *et al.* Machine learning models for streamflow regionalization in a tropical watershed. *Journal of Environmental Management*, v. 280, p. 111713, 15 fev. 2021.

FRIEDMAN, J. H. Greedy function approximation: a gradient boosting machine. *The Annals of Statistics*, v. 29, n. 5, p. 1189–1232, 2001.

GALLEAR, Joseph W. *et al.* Evaluation of machine learning approaches for large-scale agricultural drought forecasts to improve monitoring and preparedness in Brazil. *Natural Hazards and Earth System Sciences*, v. 25, n. 4, p. 1521–1541, 25 abr. 2025.

GHOBADI, F.; KANG, D. *Application of Machine Learning in Water Resources Management: A Systematic Literature Review*. *Water*, v. 15, n. 04, 2023.

JAMES, G. *et al.* *An Introduction to Statistical Learning: with Applications in Python*. Cham: Springer, 2023.

KOCHHAR, Aarti *et al.* Prediction and forecast of pre-monsoon and post-monsoon groundwater level: using deep learning and statistical modelling. *Modeling Earth Systems and Environment*, v. 8, n. 2, p. 2317–2329, 1 jun. 2022.

KUHN, M.; JOHNSON, K. *Feature Engineering and Selection: A Practical Approach for Predictive Models*. Boca Raton: CRC Press, 2019.

MORETTIN, P. A.; SINGER, J. da M. *Estatística e ciência de dados*. Rio de Janeiro: LTC, 2025.

MUNICH RE. *Natural Catastrophes 2014. Analyses, Assessments, Positions. TOPICS-GEO 2014*. Munich: [s.n.], 2015.

NOBRE, Carlos A. *et al.* Some Characteristics and Impacts of the Drought and Water Crisis in Southeastern Brazil during 2014 and 2015. *Journal of Water Resource and Protection*, v. 8, n. 2, p. 252–262, 4 fev. 2016.

SANTOS, Celso; MORAIS, Bruno; SILVA, Gustavo. Drought forecast using Artificial Neural Network for three hydrological zones in San Francisco river basin. *IAHS Publication*, v. 333, p. 302–312, 1 jan. 2009.

TEIMOORI, Sadaf; OLYA, Mohammad Hessam; MILLER, Carol J. Groundwater level monitoring network design with machine learning methods. *Journal of Hydrology*, v. 625, p. 130145, 1 out. 2023.

ZHANG, B. *et al.* Explainable machine learning for the prediction and assessment of complex drought impacts. *Science of the Total Environment*, v. 898, p. 165509, 10 nov. 2023.

AGRADECIMENTOS

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior - Brasil (CAPES) - Código de Financiamento 001. Os autores gostariam de agradecer à CAPES, à Fundação Centro Tecnológico de Hidráulica (FCTH) e ao Programa de Pós-Graduação em Engenharia Civil (PPGEC) da Escola Politécnica da Universidade de São Paulo, por apoiarem a realização deste trabalho.