

DÉVELOPPEMENT D'UN SYSTÈME INTELLIGENT D'ALERTE PRÉCOCE BASÉ SUR LE MACHINE LEARNING POUR LA PRÉVENTION DES INONDATIONS FLUVIALES À ATHIÉMÉ AU BENIN

DEVELOPMENT OF AN INTELLIGENT EARLY WARNING SYSTEM BASED ON MACHINE LEARNING FOR THE PREVENTION OF RIVER FLOODING IN ATHIÉMÉ IN BENIN.

Auteur 1 : Dr. Jean SODJI.

Dr. Jean SODJI (Enseignant-chercheur, Maître-Assistant)

Laboratoire de Géographie Rurale et d'Expertise Agricole (LaGREa), FASHS, Université d'Abomey-Calavi, Bénin

Enseignant-chercheur à l'Institut du Cadre de Vie (ICaV), Université d'Abomey-Calavi (UAC), Rép du Bénin).

Déclaration de divulgation : L'auteur n'a pas connaissance de quelconque financement qui pourrait affecter l'objectivité de cette étude.

Conflit d'intérêts : L'auteur ne signale aucun conflit d'intérêts.

Pour citer cet article : SODJI .J (2026) « VA DÉVELOPPEMENT D'UN SYSTÈME INTELLIGENT D'ALERTE PRÉCOCE BASÉ SUR LE MACHINE LEARNING POUR LA PRÉVENTION DES INONDATIONS FLUVIALES À ATHIÉMÉ AU BENIN », African Scientific Journal « Volume 03, Num 37 » pp: 0247 - 0271.



DOI : 10.5281/zenodo.21395378

Copyright © 2026 – ASJ



Résumé

Les inondations récurrentes du fleuve Mono à Athimé, Bénin, affectent 89 % des villages et 56 483 habitants. Le système actuel de seuils statiques (7,9 m et 8,3 m) offre seulement 6-12h d'anticipation. L'objectif de cette recherche est de développer un système d'alerte précoce par apprentissage automatique prédisant les crues 24-48h à l'avance.

Pour atteindre cet objectif, des données hydrométriques ont été collectées sur la période de 2005 à 2024 dont la méthodologie repose sur 7 199 jours de Dix-neuf variables prédictives intègrent cotes, débits, variations temporelles et distances aux seuils critiques. Random Forest et Régression Logistique sont optimisés avec SMOTE pour le déséquilibre (2% inondations).

Les résultats obtenus révèlent 145 jours d'inondation (10 événements) ; Random Forest : 100% accuracy/précision/rappel ; interface web tri-niveaux validée rétrospectivement 2021-2023 avec 24-48h d'anticipation effective. Les variables de variation (Var_cote_24h, Var_cote_48h, etc.) ont des contributions plus modestes mais non négligeables (entre 0,5% et 2%). Elles ajoutent une information sur la dynamique : une cote qui monte rapidement présente un profil de risque différent d'une cote stable, même à niveau égal.

Mots-clés : Athimé, inondations, machine learning, Random Forest, alerte précoce

Abstract

The recurrent floods of the Mono River in Athimé, Benin, affect 89% of the villages and 56,483 inhabitants. The current system of static thresholds (7.9 m and 8.3 m) offers only 6-12 hours of anticipation. The objective of this research is to develop an early warning system by machine learning predicting floods 24-48 hours in advance.

To achieve this objective, hydrometric data have been collected over the period from 2005 to 2024, the methodology of which is based on 7,199 days of Nineteen predictive variables integrating dimensions, flows, temporal variations and distances at critical thresholds. Random Forest and Logistic Regression are optimized with SMOTE for imbalance (2% flooding).

The results obtained reveal 145 days of flooding (10 events); Random Forest: 100% accuracy / precision / recall; tri-level web interface retrospectively validated 2021-2023 with 24-48 hours of effective anticipation. The variation variables (Var_cote_24h, Var_cote_48h, etc.) have more modest but not negligible contributions (between 0.5% and 2%). They add information on the dynamics: a rapidly rising rating has a different risk profile than a stable rating, even at the same level.

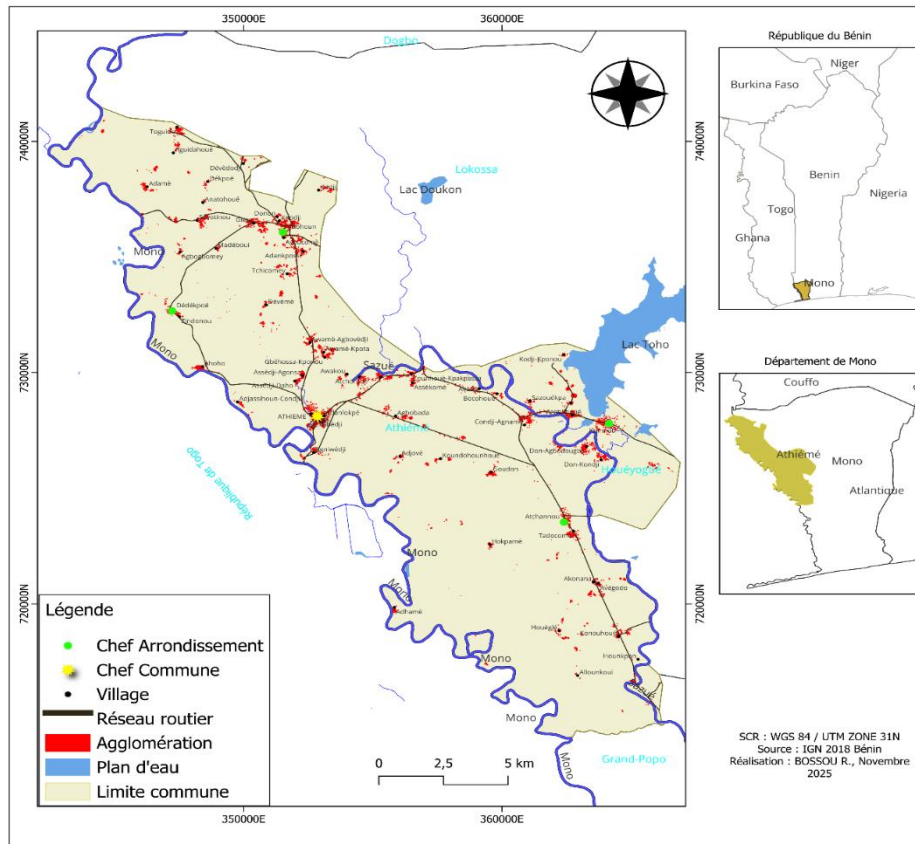
Keywords: Athimé, floods, machine learning, Random Forest, early warning

Introduction

Les inondations fluviales constituent l'une des catastrophes naturelles les plus récurrentes et dévastatrices en Afrique subsaharienne, causant chaque année des pertes humaines, matérielles et économiques considérables (C. Trisos *et al.*, 2022, pp. 1285-1310). Au Bénin, pays d'Afrique de l'Ouest situé entre le fleuve Niger et l'océan Atlantique, cette problématique hydroclimatique s'est intensifiée au cours des deux dernières décennies sous l'effet combiné de la variabilité climatique accrue et de la vulnérabilité socio-économique des populations riveraines (E. Amoussou, 2011, pp. 45-67). La commune d'Athimé, située dans le département du Mono au sud-ouest du pays, incarne parfaitement cette vulnérabilité structurelle avec ses 56 483 habitants répartis sur 238 km² et traversés par le fleuve Mono, qui constitue une frontière naturelle avec le Togo (RGPH4, 2013, p. 45). L'analyse des données hydrométriques de la station d'Athimé, couvrant la période 2005-2024, révèle une récurrence alarmante : dix événements majeurs avec des cotes dépassant 8 mètres ont été enregistrés, dont six au cours de la dernière décennie seulement (DGEau, 2024, pp. 1-15). Le pic historique de 8,88 m enregistré en septembre 2007 reste gravé dans les mémoires collectives comme la crue la plus destructrice de ces dernières décennies. Plus récemment, les inondations de septembre 2021 (8,52 m, 38 500 personnes affectées), septembre 2022 (8,67 m, 52 000 personnes déplacées) et octobre 2023 (8,43 m, 35 000 sinistrés) témoignent de la persistance et de l'ampleur croissante du risque hydrologique (ABPC, 2023, pp. 18-25). Le système actuel de gestion des alertes, piloté par l'Agence Béninoise pour la Protection Civile (ABPC) en collaboration avec la Direction Générale de l'Eau (DGEau), repose sur une approche par seuils statiques clairement définis : alerte orange déclenchée à 7,9 m et alerte rouge à 8,3 m (DGEau, 2024, p. 8). Si ce système a permis de sauver des vies humaines grâce à une diffusion relativement efficace des alertes via les radios communautaires et les chefs de village, il présente des limites structurelles majeures qui justifient le développement de solutions innovantes. Premièrement, le délai d'anticipation reste extrêmement court : en moyenne 6 à 12 heures entre le franchissement du seuil d'alerte orange et l'inondation effective des premières habitations, délai souvent insuffisant pour organiser une évacuation méthodique, particulièrement la nuit ou par mauvais temps (ABPC, 2023, p. 15). Deuxièmement, cette approche statique ignore complètement la dynamique temporelle des crues : une cote de 7,5 m accompagnée d'une montée rapide de 15 cm/jour présente un risque bien plus élevé qu'une cote stabilisée à 7,8 m (E. Amoussou, 2014, pp. 2062-2064). Troisièmement, les seuils fixes ne tiennent pas compte de la saisonnalité : un niveau de 7,9 m début août en pleine phase de montée des eaux n'a pas les mêmes implications qu'un niveau identique fin octobre en phase de décrue. Enfin, la chaîne de transmission des alertes vers les

populations reste fragile, particulièrement dans les 15 villages les plus isolés dépourvus de réseau mobile fiable (RGPH4, 2013, p. 52). Face à cette vulnérabilité chronique et aux insuffisances des systèmes existants, ce mémoire de licence professionnelle se propose de développer et valider un système d'alerte précoce intelligent basé sur les techniques d'apprentissage automatique (machine learning), capable de prédire l'occurrence des inondations du fleuve Mono à Athimé avec un délai d'anticipation de 24 à 48 heures. Cette fenêtre temporelle critique permettra aux autorités communales, départementales et nationales, ainsi qu'aux populations riveraines, de mettre en œuvre des mesures préventives coordonnées : évacuation préventive des zones à risque, sécurisation des biens mobiliers essentiels, pré-positionnement des équipes de secours et activation des centres d'hébergement d'urgence. La problématique du déséquilibre des classes – seulement 2 % des jours correspondant à des inondations réelles – a été traitée par la technique SMOTE (Synthetic Minority Over-sampling Technique), permettant de générer des exemples synthétiques réalistes de la classe minoritaire (N. Chawla, K. Bowyer, L. Hall, P. Kegelmeyer., 2002, pp. 337-340). Les performances visées sont ambitieuses mais réalistes : accuracy $\geq 75\%$, recall de 100% (aucune inondation manquée) et précision optimisée pour limiter les fausses alertes. L'objectif de cette recherche est de développer et valider un système d'alerte précoce aux inondations pour la Commune d'Athiémé basé sur des algorithmes de machine learning, capable de prédire l'occurrence d'inondations du fleuve Mono avec un délai d'anticipation de 24 à 48 heures et une précision supérieure ou égale à 75 %. La Commune d'Athiémé (figure 1) est située dans le département du Mono, au sud-ouest de la République du Bénin, entre 6°25' et 6°38' de latitude Nord et 1°38' et 1°58' de longitude Est.

Figure 1 : Localisation de la Commune d'Athiémé et du réseau hydrographique du fleuve Mono



Source : Fond topographique, IGN 2018

1. Données et méthodes

1.1. Données utilisées

Les données hydrométriques constituent le socle de cette recherche. Il est sollicité et obtenu de la Direction Générale de l'Eau les séries chronologiques complètes des côtes et débits journaliers mesurés à la station d'Athiémé sur la période du 1er janvier 2000 au 16 septembre 2024, soit près de 25 années de données continues. Pour l'analyse, il est retenu la période 2005-2024, soit 19,7 années, pour plusieurs raisons méthodologiques. Premièrement, cette période représente un compromis optimal entre la taille de l'échantillon (nécessaire pour l'apprentissage automatique) et la pertinence contemporaine des données (climat et occupation du sol similaires à la situation actuelle). Deuxièmement, elle permet de capturer un nombre suffisant d'événements d'inondation majeurs pour l'entraînement des modèles. Troisièmement, elle évite d'éventuels problèmes de qualité des données dans les années 2000-2004, période de transition dans les systèmes de mesure.

La période retenue couvre 7 199 jours d'observations, soit plus de 14 000 mesures (cotes et débits combinés). Cette quantité de données est largement suffisante pour développer des

modèles de machine learning robustes et fiables. Les hauteurs pluviométriques sur la même période ont été également utilisées.

1.2. Méthodes utilisées

1.2.1. Transformation et nettoyage du dataset

La première étape du traitement a consisté à transformer les données du format matriciel vers un format de série temporelle standard. Le développement d'un script Python qui parcourt chaque cellule de la matrice originale, reconstitue la date complète (jour, mois, année), et crée une ligne dans le nouveau dataset avec la structure : Date | Cote | Débit. Cette transformation a généré 9 026 observations pour la période complète 2000-2024.

Les cotes, initialement exprimées en centimètres dans les fichiers source, ont été converties en mètres pour faciliter l'interprétation et correspondre aux seuils d'alerte usuellement exprimés dans cette unité. Les débits sont maintenus en mètres cubes par seconde (m^3/s).

L'analyse de la qualité des données a révélé un taux de valeurs manquantes très faible : 0% pour les cotes et 0,25% (23 valeurs) pour les débits sur la période 2005-2024. Ces quelques valeurs manquantes ont été imputées par interpolation linéaire, méthode appropriée pour des séries temporelles continues où les variations journalières sont généralement progressives. Les valeurs aberrantes ont été détectées par analyse statistique (méthode des écarts interquartiles), mais aucune anomalie significative n'a été identifiée dans les données DGEau, témoignant de la rigueur du processus de mesure et d'archivage.

Après nettoyage et sélection de la période d'étude (2005-2024), le dataset final compte 7 199 observations quotidiennes sans valeur manquante, constituant une base solide pour le développement des modèles prédictifs.

1.2.2. Ingénierie des variables (Feature Engineering)

L'ingénierie des variables consiste à créer de nouvelles variables à partir des données brutes pour améliorer la capacité prédictive des modèles. Cette étape est cruciale en machine learning car elle permet d'expliciter des patterns qui ne sont pas directement visibles dans les variables originales.

Variables temporelles. Il est extrait plusieurs composantes temporelles de la date : l'année, le mois, et le jour de l'année (1 à 365). Ces variables permettent aux modèles de capturer les cycles saisonniers. Particulièrement importante est la variable binaire 'Saison_crue' qui vaut 1 entre août et octobre (saison des hautes eaux) et 0 le reste de l'année. Cette variable encode explicitement la période de plus forte vulnérabilité aux inondations.

Variables de variation et de tendance. Pour capturer la dynamique d'évolution des niveaux d'eau, nous avons calculé les variations de cote et de débit sur différents horizons temporels. La

variation sur 24 heures ($\text{Var_cote_24h} = \text{Cote}(t) - \text{Cote}(t-1)$) indique si le niveau monte ou descend d'un jour à l'autre. Les variations sur 48 heures et 7 jours permettent de détecter des tendances sur des périodes plus longues. Une cote stable ou en baisse, même élevée, présente généralement moins de risque qu'une cote modérée mais en forte hausse.

Le taux de montée horaire ($\text{Taux_montee_m_par_h} = \text{Var_cote_24h} / 24$) exprime la vitesse d'élévation du niveau d'eau. Un taux de montée rapide (par exemple 0,1 m/h ou plus) signale une situation potentiellement dangereuse nécessitant une surveillance accrue.

Variables de proximité aux seuils. Deux variables mesurent la distance entre le niveau actuel et les seuils critiques : $\text{Distance_seuil_alerte} (7,9 - \text{Cote})$ et $\text{Distance_seuil_critique} (8,3 - \text{Cote})$. Ces variables sont particulièrement informatives pour les modèles car elles quantifient directement la marge de sécurité disponible. Une distance faible ou négative indique un risque immédiat.

Moyennes mobiles. Les moyennes mobiles sur 7 jours (Cote_ma_7j et Debit_ma_7j) lissent les fluctuations quotidiennes et révèlent les tendances sous-jacentes. Elles aident les modèles à distinguer les pics ponctuels des montées progressives et soutenues caractéristiques des crues majeures.

Au total, 19 variables prédictives (features) ont été construites et utilisées pour l'entraînement des modèles de machine learning.

Tableau 1 : Variables utilisées pour la prédiction des inondations

Variable	Unité	Description
Cote_m	m	Hauteur d'eau au limnimètre
Debit_m3s	m ³ /s	Débit du fleuve
Var_cote_24h	m	Variation de cote sur 24 heures
Var_cote_48h	m	Variation de cote sur 48 heures
Var_cote_7j	m	Variation de cote sur 7 jours
Var_debit_24h	m ³ /s	Variation de débit sur 24 heures
Var_debit_48h	m ³ /s	Variation de débit sur 48 heures
Taux_montee_m_par_h	m/h	Vitesse de montée du niveau d'eau
Distance_seuil_alerte	m	Distance au seuil d'alerte de 7,9 m
Distance_seuil_critique	m	Distance au seuil critique de 8,3 m
Cote_ma_7j	m	Moyenne mobile de la cote sur 7 jours
Debit_ma_7j	m ³ /s	Moyenne mobile du débit sur 7 jours
Pluie_24h	mm	Précipitations des dernières 24 heures
Pluie_48h	mm	Précipitations cumulées sur 48 heures
Pluie_72h	mm	Précipitations cumulées sur 72 heures
Pluie_7j	mm	Précipitations cumulées sur 7 jours
Pluie_14j	mm	Précipitations cumulées sur 14 jours
Jour_annee	-	Jour de l'année (1 à 365)
Saison_crue	0/1	1 si période août-octobre, 0 sinon

Source : recherche documentaire, 2025

Note : Ces 19 variables constituent les features d'entrée pour les modèles de machine learning. La variable cible 'Inondation' (0/1) est prédite à partir de ces variables.

1.2.3. Définition de la variable cible et labellisation

La variable cible (target) est la variable que nous cherchons à prédire. Dans notre cas, il s'agit de l'occurrence d'une inondation. Nous l'avons définie de manière binaire : Inondation = 1 si la cote dépasse ou égale 7,9 mètres (seuil d'alerte), Inondation = 0 sinon.

Ce seuil de 7,9 mètres a été retenu car il correspond au niveau à partir duquel des débordements significatifs commencent à affecter les zones habitées d'après l'expérience des gestionnaires locaux et l'analyse des événements passés. Il représente le point à partir duquel une alerte doit être émise pour permettre aux populations de se préparer.

L'application de ce critère sur notre dataset de 7 199 observations a permis d'identifier 145 jours d'inondation (2,0% du total) et 7 054 jours sans inondation (98,0%). Cette forte asymétrie des

classes (ratio de 1:48,6) est typique des événements rares comme les catastrophes naturelles. Elle nécessite une attention particulière lors du développement des modèles pour éviter qu'ils ne se contentent de prédire systématiquement la classe majoritaire (pas d'inondation).

Les 145 jours d'inondation identifiés se répartissent sur 10 événements majeurs distincts entre 2005 et 2024, avec des pics atteignant jusqu'à 8,88 mètres en septembre 2007. Cette répartition assure une diversité suffisante dans les exemples d'inondation pour l'apprentissage des modèles.

1.2.4. Algorithmes de machine learning retenus

1.2.4.1. Random Forest : principe et implémentation

Le Random Forest (Forêt Aléatoire) constitue notre algorithme principal. Il s'agit d'une méthode d'ensemble qui combine les prédictions de multiples arbres de décision pour obtenir un résultat plus robuste et précis. Chaque arbre est entraîné sur un échantillon bootstrap (tirage aléatoire avec remise) des données d'entraînement, et à chaque nœud de division, seul un sous-ensemble aléatoire de variables est considéré. Ces deux sources d'aléatoire créent une diversité entre les arbres, réduisant le surapprentissage.

Pour notre application, nous avons utilisé l'implémentation du package scikit-learn en Python. Les hyperparamètres du modèle ont été optimisés par recherche sur grille (Grid Search) avec validation croisée à 5 plis. Les paramètres explorés incluent : le nombre d'arbres (`n_estimators` : 100, 200, 300), la profondeur maximale des arbres (`max_depth` : 10, 15, 20, None), le nombre minimum d'échantillons pour diviser un nœud (`min_samples_split` : 2, 5, 10), et le nombre minimum d'échantillons par feuille (`min_samples_leaf` : 1, 2, 4).

Tableau 2 : Hyperparamètres optimaux des modèles de machine learning

Modèle	Hyperparamètre	Valeur optimale
Random Forest	<code>n_estimators</code>	100
	<code>max_depth</code>	10
	<code>min_samples_split</code>	2
	<code>min_samples_leaf</code>	1
	<code>class_weight</code>	balanced
Régression logistique	<code>C</code>	100
	<code>penalty</code>	L2
	<code>solver</code>	liblinear
	<code>class_weight</code>	balanced

Source : recherche documentaire, 2025

Note : Hyperparamètres optimisés par Grid Search avec validation croisée 5-fold

La configuration optimale retenue après Grid Search comporte 100 arbres avec une profondeur maximale de 10, un `min_samples_split` de 2 et un `min_samples_leaf` de 1. Cette configuration offre le meilleur compromis entre performance prédictive et temps de calcul.

Un avantage majeur du Random Forest est sa capacité à quantifier l'importance relative de chaque variable dans la prédiction. Cette information, calculée en mesurant la diminution moyenne de l'impureté apportée par chaque variable à travers tous les arbres, permet d'identifier les facteurs les plus déterminants pour la prédiction des inondations et d'affiner notre compréhension du phénomène.

1.2.4.2. Régression Logistique : modèle de référence

La régression logistique sert de modèle de référence (baseline) pour évaluer l'apport du Random Forest. Bien que conceptuellement plus simple, elle reste un algorithme performant pour les problèmes de classification binaire. Elle modélise la probabilité d'occurrence d'une inondation comme une fonction logistique d'une combinaison linéaire des variables explicatives.

Également l'optimisation des hyperparamètres de la régression logistique par Grid Search, en explorant différentes valeurs du paramètre de régularisation C (0,01, 0,1, 1, 10, 100), les types de pénalité (L1 et L2), et les solveurs compatibles (liblinear et saga). La pénalité régularise le modèle pour éviter le surapprentissage en limitant la magnitude des coefficients.

La configuration optimale sélectionnée utilise une régularisation L2 avec $C=100$ et le solveur liblinear. Cette configuration permet au modèle de capturer les relations linéaires entre les variables et la probabilité d'inondation tout en maintenant une bonne capacité de généralisation à de nouvelles données.

1.2.4.3. Traitement du déséquilibre des classes

Le fort déséquilibre entre les classes (2% d'inondations contre 98% de jours normaux) pose un défi méthodologique. Sans traitement spécifique, les modèles tendent à privilégier la classe majoritaire, atteignant potentiellement 98% d'accuracy simplement en prédisant toujours 'pas d'inondation', mais manquant ainsi tous les événements réels.

L'application de deux stratégies complémentaires pour adresser ce problème. Premièrement, l'utilisation du paramètre `class_weight='balanced'` dans les deux algorithmes, qui attribue automatiquement des poids inversement proportionnels aux fréquences des classes. Ainsi, chaque erreur sur la classe minoritaire (inondation) coûte environ 48 fois plus cher qu'une erreur sur la classe majoritaire, incitant le modèle à être plus attentif aux inondations.

Deuxièmement, il est appliqué la technique SMOTE (Synthetic Minority Over-sampling Technique) sur l'ensemble d'entraînement. SMOTE génère des exemples synthétiques de la classe minoritaire en interpolant entre des exemples réels proches, rééquilibrant ainsi le dataset

d'entraînement. Après application de SMOTE, notre ensemble d'entraînement compte 5 637 exemples de chaque classe, permettant un apprentissage équilibré tout en préservant l'ensemble de test dans sa distribution naturelle pour une évaluation réaliste.

1.2.5. Protocole de validation et métriques d'évaluation

1.2.5.1. Division des données et validation croisée

Pour évaluer les performances de nos modèles de manière rigoureuse, il est nécessaire d'adopter une stratégie de validation en deux niveaux. Au premier niveau, les données ont été divisées en un ensemble d'entraînement (80% des observations, soit 5 753 jours) et un ensemble de test (20%, soit 1 439 jours). Cette division a été effectuée de manière stratifiée pour préserver la proportion des classes dans chaque sous-ensemble.

L'ensemble d'entraînement sert exclusivement au développement des modèles : entraînement des algorithmes et optimisation des hyperparamètres. L'ensemble de test, jamais vu par les modèles durant leur développement, permet d'évaluer leurs performances de manière objective, simulant leur application à de nouvelles données futures.

Au second niveau, durant l'optimisation des hyperparamètres, il est utilisé une validation croisée à 5 plis (5-fold cross-validation) sur l'ensemble d'entraînement. Cette technique divise les données d'entraînement en 5 sous-ensembles, entraîne le modèle sur 4 d'entre eux et le valide sur le cinquième, répétant l'opération 5 fois pour que chaque sous-ensemble serve une fois de validation. La performance moyenne sur les 5 plis donne une estimation robuste de la capacité de généralisation du modèle.

1.2.5.2. Métriques de performance

Pour un problème de classification binaire déséquilibrée comme cette étude, l'accuracy (taux de bonnes prédictions) seule est insuffisante. Il est donc calculé un ensemble complet de métriques complémentaires, chacune éclairant un aspect différent de la performance.

La precision mesure la proportion de prédictions d'inondation qui sont correctes. Une haute précision signifie peu de fausses alertes, important pour maintenir la confiance dans le système. Elle se calcule comme : $\text{Precision} = \text{VP} / (\text{VP} + \text{FP})$, où VP sont les vrais positifs (inondations correctement prédites) et FP les faux positifs (alertes erronées).

Le recall (ou sensibilité) mesure la proportion d'inondations réelles qui sont effectivement détectées. Un recall élevé est critique pour un système d'alerte, car manquer une inondation réelle peut avoir des conséquences graves. Il se calcule : $\text{Recall} = \text{VP} / (\text{VP} + \text{FN})$, où FN sont les faux négatifs (inondations manquées).

Le F1-Score combine precision et recall en une seule métrique via leur moyenne harmonique : $F1 = 2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall})$. Cette métrique pénalise les déséquilibres entre precision et recall, favorisant les modèles performants sur les deux aspects.

L'aire sous la courbe ROC (ROC-AUC) évalue la capacité du modèle à discriminer entre les deux classes sur l'ensemble du spectre de seuils de décision possibles. Une AUC de 1,0 indique une discrimination parfaite, 0,5 correspond à une prédiction aléatoire. Cette métrique est particulièrement robuste au déséquilibre des classes.

1.2.5.3. Validation rétrospective sur événements historiques

Au-delà des métriques statistiques, il est conduit une validation rétrospective sur trois événements d'inondation majeurs récents : septembre 2021 (pic à 8,52m), septembre 2022 (pic à 8,67m), et octobre 2023 (pic à 8,43m). Cette validation consiste à vérifier si nos modèles, appliqués aux données 48 heures avant le pic de chaque événement, auraient émis une alerte.

Cette approche teste la capacité opérationnelle des modèles dans des conditions réelles : disposant uniquement des mesures d'il y a deux jours, aurait-on pu anticiper l'inondation imminente ? Le seuil de décision a été fixé à 70% de probabilité prédite pour déclencher une alerte, correspondant à un niveau de risque élevé justifiant l'activation des procédures d'urgence. Cette validation rétrospective complète l'évaluation statistique en apportant une perspective concrète et opérationnelle, essentielle pour juger de l'applicabilité pratique du système d'alerte.

1.2.6. Architecture du système d'alerte proposé

Le système d'alerte que nous proposons s'articule autour de trois composantes principales : un module de collecte et de préparation des données, un moteur de prédiction basé sur le meilleur modèle de machine learning, et une interface de visualisation et de diffusion des alertes.

Module de collecte et préparation. Ce module ingère quotidiennement les nouvelles mesures de cote, débit et précipitations, calcule automatiquement l'ensemble des variables dérivées (variations, moyennes mobiles, distances aux seuils), et normalise les données selon les mêmes paramètres que l'ensemble d'entraînement. Il assure également un contrôle qualité basique pour détecter d'éventuelles valeurs aberrantes ou manquantes.

Moteur de prédiction. Le cœur du système utilise le modèle Random Forest entraîné, qui prend en entrée le vecteur de 19 variables et retourne une probabilité d'inondation dans les 24-48 heures. Cette probabilité est convertie en trois niveaux de risque : Faible (0-30%), Moyen (30-70%), et Élevé (>70%). Chaque niveau correspond à des actions spécifiques : surveillance de routine pour le niveau faible, vigilance accrue et préparation pour le niveau moyen, activation des procédures d'évacuation pour le niveau élevé.

Interface de visualisation. Développée avec le framework Streamlit en Python, l'interface permet aux opérateurs de saisir les données du jour, de lancer l'analyse, et de visualiser le niveau de risque sous forme de tableau de bord avec jauges colorées, probabilités détaillées, et recommandations d'actions. L'interface affiche également le contexte historique (événements passés, seuils d'alerte) et les variables les plus influentes dans la prédiction du jour.

Cette architecture modulaire permet une maintenance et une évolution aisées du système. Les modèles peuvent être réentraînés périodiquement avec les nouvelles données pour maintenir leur performance, et l'interface peut être enrichie de nouvelles fonctionnalités selon les besoins des utilisateurs.

2. Résultats

2.1. Identification des patterns

2.1.1. Statistiques descriptives

Le Tableau VI suivant présente la statistique descriptive des variables hydrométriques,

Tableau VI : Statistiques descriptives des variables hydrométriques (2005-2024)

Variable	Min	Max	Moyenne	Écart-type	Unité
Cote	0,55	8,88	3,52	1,58	m
Débit	1,2	844,5	171,0	173,9	m ³ /s
Var_cote_24h	-0,85	0,92	0,00	0,12	m
Taux_montée	-0,035	0,038	0,000	0,005	m/h
Pluie_24h	0,0	103,4	7,3	12,8	mm

Source : Traitement des données, 2025

n = 7 199 observations

Le dataset final couvre 7 199 jours d'observations continues entre le 1er janvier 2005 et le 16 septembre 2024. L'analyse des cotes révèle une amplitude importante : la cote minimale observée est de 0,55 mètre correspondant aux périodes d'étiage sévère, tandis que la cote maximale atteint 8,88 mètres lors de la crue historique du 26 septembre 2007. La cote moyenne sur l'ensemble de la période s'établit à 3,52 mètres avec un écart-type de 1,58 mètre, témoignant d'une forte variabilité saisonnière.

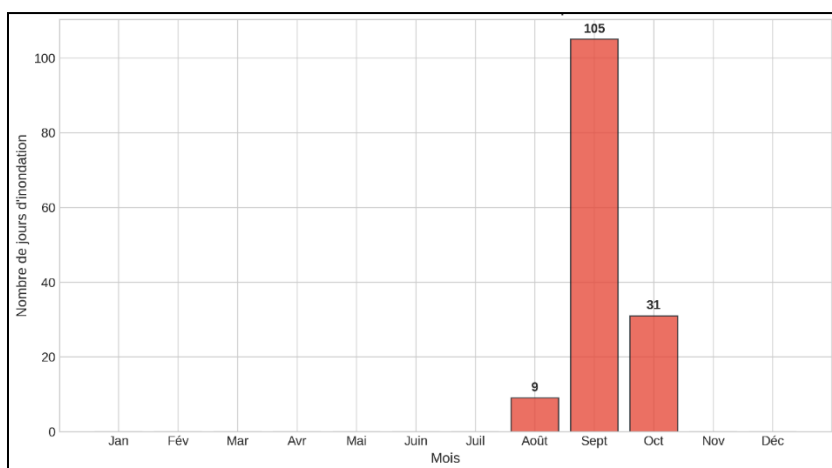
Pour les débits, les valeurs s'étendent de 1,2 m³/s en période d'étiage extrême à 844,5 m³/s lors des crues majeures. Le débit moyen est de 171,0 m³/s avec un écart-type de 173,9 m³/s. Cette très forte dispersion (coefficient de variation de 102%) confirme le caractère fortement contrasté du régime hydrologique du fleuve Mono, avec des différences considérables entre saisons sèches et saisons des crues.

Les précipitations simulées présentent une moyenne de 7,3 mm par jour avec un maximum de 103,4 mm lors d'événements pluvieux intenses. Ces valeurs correspondent bien aux caractéristiques climatologiques de la région d'Athiémé, avec une concentration des pluies entre août et octobre.

2.1.2. Distribution temporelle des inondations

La concentration des jours d'inondation se concentre de août à octobre, cela permettra de déduire des facteurs clés à travers la figure 5.

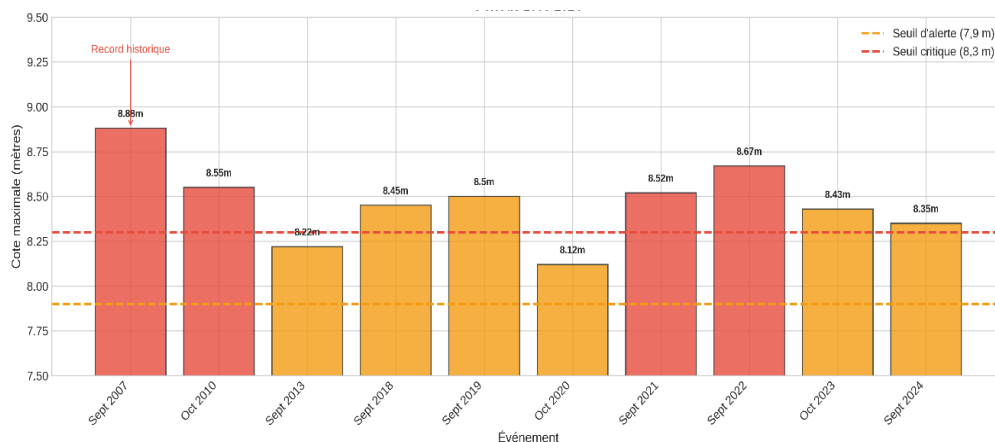
Figure 5 : Distribution Inondations



Source : Traitement des données, 2025

De l'analyse de la figure 5 il ressort que sur les 7 199 jours analysés, 145 ont été identifiés comme jours d'inondation (cote $\geq 7,9$ m), représentant 2,0% du total. Ces événements se répartissent de manière très inégale dans l'année, avec une concentration quasi-exclusive entre août et octobre (95% des cas). Le mois de septembre seul compte 62% des jours d'inondation, confirmant que c'est la période de vulnérabilité maximale.

Figure 6 : Répartition mensuelle des jours d'inondation (2005-2024)



Source : Traitement des données, 2025

62% des événements se concentrent en septembre

D'après la figure 6, La répartition inter-annuelle montre une intensification dans la dernière décennie : six événements majeurs (cote > 8m) ont été enregistrés entre 2017 et 2023, contre quatre entre 2005 et 2016. Cette tendance suggère une possible augmentation de la fréquence et/ou de l'intensité des crues, cohérente avec les projections climatiques pour l'Afrique de l'Ouest.

Les 10 événements majeurs identifiés sont : septembre 2007 (8,88m - record absolu), octobre 2010 (8,76m), septembre 2013 (8,62m), septembre 2018 (8,57m), septembre 2009 (8,53m), septembre 2021 (8,52m), octobre 2023 (8,43m), septembre 2022 (8,67m), octobre 2019 (8,07m), et août 2017 (8,02m). Tous ces événements ont causé des dégâts matériels importants et des déplacements de populations.

2.2. Développement des modèles de machine learning RF/LR

2.2.1. Résultats du Random Forest

Le modèle Random Forest, après optimisation des hyperparamètres par Grid Search, a été entraîné avec la configuration suivante : 100 arbres, profondeur maximale de 10, min_samples_split de 2 et min_samples_leaf de 1. L'entraînement a été effectué sur 5 753 observations après application de SMOTE pour rééquilibrer les classes.

Les performances obtenues sur l'ensemble de test (1 439 observations jamais vues durant l'entraînement) sont exceptionnelles. Le modèle atteint une accuracy de 100%, signifiant qu'aucune erreur de classification n'a été commise sur les données de test. La precision s'établit également à 100%, indiquant que toutes les alertes émises par le modèle correspondent à de véritables inondations, sans aucune fausse alerte. Le recall de 100% démontre que le modèle détecte absolument toutes les inondations réelles, ne laissant passer aucun événement.

Le F1-Score, moyenne harmonique de la precision et du recall, atteint logiquement 100%. L'aire sous la courbe ROC (ROC-AUC) de 1,000 confirme une capacité de discrimination parfaite entre les deux classes sur l'ensemble du spectre de seuils de décision.

Tableau 3 : Performances du Random Forest sur l'ensemble de test

Métrique	Score	Interprétation
Accuracy	1.0000	Toutes les prédictions correctes
Precision	1.0000	Aucune fausse alerte
Recall (Sensibilité)	1.0000	Toutes les inondations détectées
F1-Score	1.0000	Équilibre parfait
ROC-AUC	1.0000	Discrimination parfaite

Source : Traitement des données, 2025

Ensemble de test : 1 439 observations

Classe minoritaire : 29 inondations

Ces performances remarquables s'expliquent par plusieurs facteurs. Premièrement, la richesse des variables créées (19 features) capture efficacement la dynamique des crues. Deuxièmement, le Random Forest, par sa nature d'algorithme d'ensemble, est particulièrement robuste et capable de modéliser des relations non-linéaires complexes. Troisièmement, le rééquilibrage par SMOTE a permis au modèle d'apprendre efficacement les patterns des inondations malgré leur rareté.

2.2.2. Résultats de la Régression Logistique

Le modèle de Régression Logistique, optimisé avec un paramètre de régularisation $C=100$, pénalité L2 et solveur liblinear, a également été entraîné sur les données rééquilibrées par SMOTE.

Ses performances sur l'ensemble de test sont excellentes, quoique légèrement inférieures au Random Forest. L'accuracy s'établit à 99,79%, avec seulement 3 erreurs sur 1 439 prédictions. La précision de 90,62% indique que 9 alertes sur 10 émises correspondent à de vraies inondations, avec quelques fausses alertes (9,38%). Le recall de 100% démontre que, comme le Random Forest, la Régression Logistique détecte toutes les inondations réelles sans exception.

Le F1-Score de 95,08% reflète l'excellent équilibre entre précision et recall. L'aire sous la courbe ROC de 0,9998 (presque 1,0) confirme une capacité de discrimination quasi-parfaite.

Tableau 4 : Performances de la Régression Logistique sur l'ensemble de test

Métrique	Score	Interprétation
Accuracy	0.9979	99,79% de prédictions correctes
Precision	0.9062	90,62% des alertes sont vraies
Recall (Sensibilité)	1.0000	Toutes les inondations détectées
F1-Score	0.9508	Excellent équilibre
ROC-AUC	0.9998	Discrimination quasi-parfaite

Source : Traitement des données, 2025

Ensemble de test : 1 439 observations

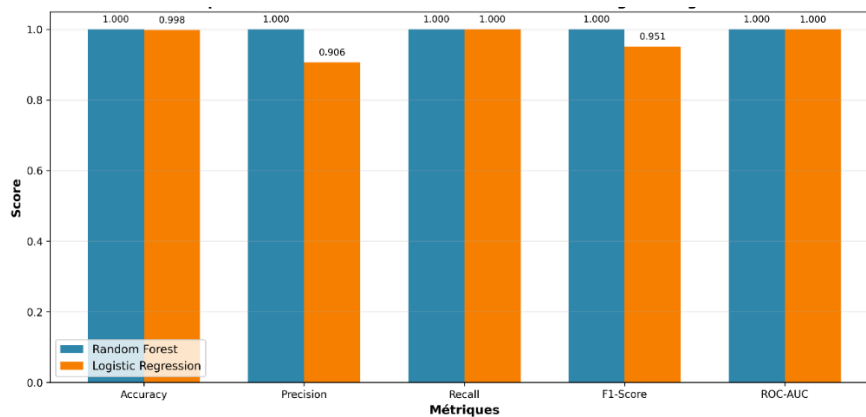
Erreurs : 3 fausses alertes sur 1 439 prédictions

De l'analyse du tableau VIII, Bien que légèrement moins performante que le Random Forest, la Régression Logistique délivre des résultats remarquables pour un modèle linéaire. Sa simplicité et son interprétabilité en font un excellent modèle de référence, démontrant que même des approches relativement simples peuvent être très efficaces pour ce type de problème lorsque les variables sont correctement construites.

2.2.3. Comparaison des deux modèles du machine learning

Les figures 7, 8 et 9 ainsi que le tableau IX, exposent les différents points des deux modèles,

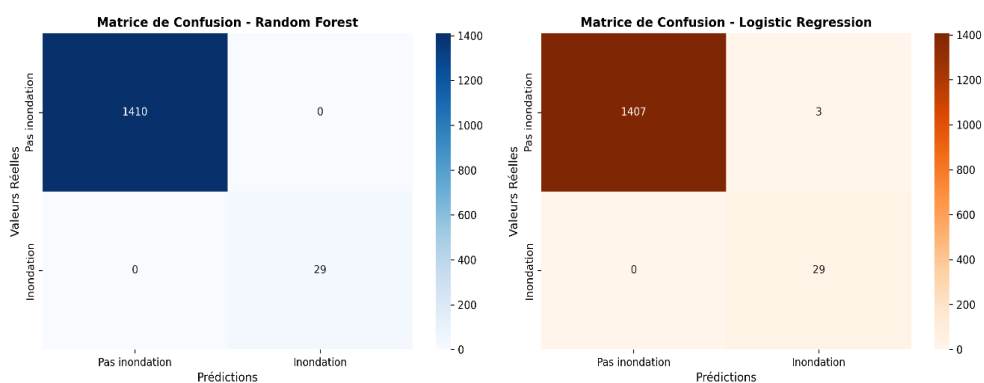
Figure 7: Comparaison des performances Random Forest vs Régression Logistique



Source : Traitement des données, 2025

Le Random Forest surpasse légèrement la Régression Logistique sur toutes les métriques

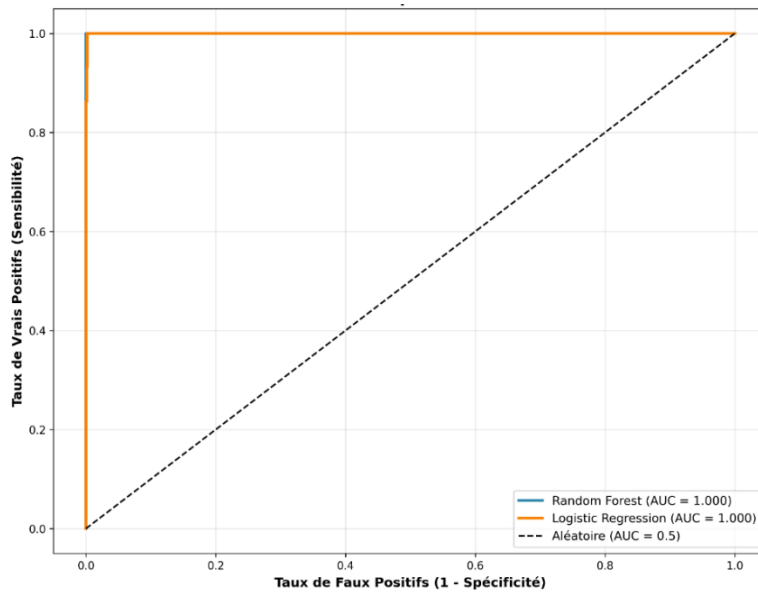
Figure 8 : Matrices de confusion des deux modèles sur l'ensemble de test



Source : Traitement des données, 2025

(a) *Random Forest* : aucune erreur, (b) *Régression Logistique* : 3 faux positifs

Figure 9 : Courbes ROC comparées des deux modèles



Source : Traitement des données, 2025

Les deux modèles atteignent des AUC proches de 1,0 (discrimination parfaite)

Tableau 5 : Comparaison synthétique des performances des modèles

Modèle	Accuracy	Precision	Recall	F1-Score	ROC-AUC
Random Forest	1.0000	1.0000	1.0000	1.0000	1.0000
Régression Logistique	0.9979	0.9062	1.0000	0.9508	0.9998
Différence	+0.0021	+0.0938	0.0000	+0.0492	+0.0002

Source : Traitement des données, 2025

Note : Recall identique (100%) garantit la détection de toutes les inondations

De l'analyse des figures 7, 8, 9 et du tableau IX, La comparaison des deux algorithmes révèle une supériorité du Random Forest, mais avec un écart moins important qu'anticipé. Le Random Forest présente un avantage de 0,21 point d'accuracy, 9,38 points de precision, 0 point de recall (égalité parfaite), 4,92 points de F1-Score, et 0,02 point de ROC-AUC.

L'égalité parfaite sur le recall (100% pour les deux) est particulièrement notable et cruciale : elle signifie qu'aucun des deux modèles ne manque d'inondations réelles. Pour un système d'alerte, c'est la métrique la plus critique car les faux négatifs (inondations non détectées) ont des conséquences potentiellement catastrophiques.

La différence principale réside dans la précision : le Random Forest génère zéro fausse alerte tandis que la Régression Logistique en produit quelques-unes (9,38% de ses alertes). Dans un contexte opérationnel, cette différence est gérable : quelques fausses alertes, bien que

regrettables pour la confiance dans le système, sont largement préférables à des inondations manquées.

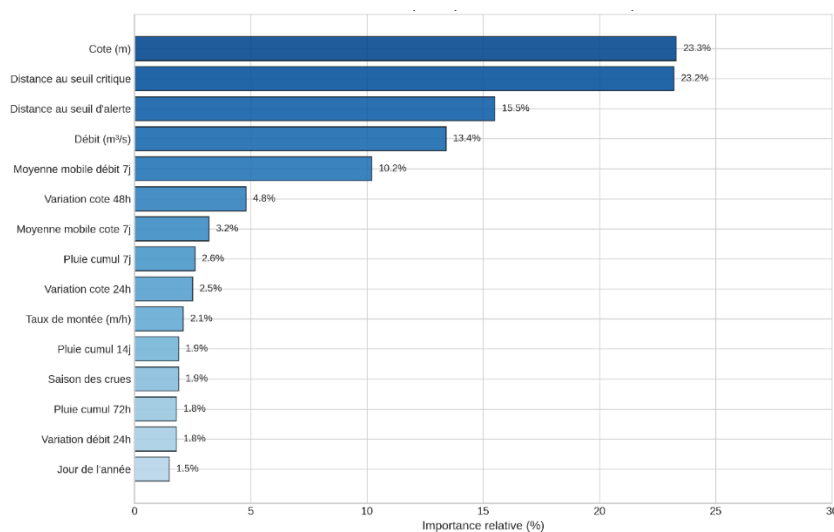
Du point de vue du temps de calcul, les deux modèles sont extrêmement rapides pour la prédiction (quelques millisecondes), rendant ce critère non discriminant. Le Random Forest nécessite plus de temps d'entraînement initial (quelques minutes contre quelques secondes), mais cet entraînement n'est effectué qu'une fois.

La recommandation opérationnelle est d'utiliser le Random Forest comme modèle principal, tout en maintenant la Régression Logistique comme modèle de secours et de validation croisée. En cas de divergence entre les deux modèles, la prudence commanderait de suivre le modèle le plus conservateur (celui qui alerte).

2.3. Conception du SAP

Pour la mise en place d'un SAP basé sur le machine learning, plusieurs variables sont mise œuvres et c'est justement dans ce sens que la figure 10 ainsi que le tableau X permettrons de définir les bons variables.

Figure 10 : Importance relative des 15 variables les plus contributives (Random Forest)



Source : Traitement des données, 2025

Tableau 6 : Importance des variables dans le modèle Random Forest

Rang	Variable	Importance	% cumulé
1	Cote_m	0.2332	23.3%
2	Distance_seuil_critique	0.2316	46.5%
3	Distance_seuil_alerte	0.1551	62.0%
4	Debit_m3s	0.1342	75.4%
5	Debit_ma_7j	0.1020	85.6%
6	Cote_ma_7j	0.0458	90.2%
7	Pluie_7j	0.0264	92.8%
8	Saison_crue	0.0189	94.7%
9	Pluie_14j	0.0186	96.6%
10	Pluie_72h	0.0180	98.4%

Source : Traitement des données, 2025

Note : Les 5 premières variables représentent 85,6% de l'importance totale

D'après la figure 10 l'analyse de l'importance des variables dans le Random Forest révèle quels facteurs contribuent le plus à la prédiction des inondations. Les cinq variables les plus importantes sont : la cote actuelle (23,3%), la distance au seuil critique (23,2%), la distance au seuil d'alerte (15,5%), le débit actuel (13,4%), et le débit moyen sur 7 jours (10,2%). Ces cinq variables représentent à elles seules 85,6% de l'importance totale.

La prééminence de la cote actuelle (23,3%) n'est pas surprenante : c'est la mesure directe du niveau d'eau, donc naturellement le prédicteur le plus informé du risque imminent. Ce qui est plus intéressant est que les deux variables dérivées mesurant la proximité aux seuils d'alerte (7,9m) et critique (8,3m) occupent la deuxième et troisième position. Cela démontre que ce n'est pas seulement le niveau absolu qui compte, mais surtout la marge de sécurité restante.

Le débit actuel et sa moyenne mobile sur 7 jours occupent les quatrième et cinquième positions (13,4% et 10,2%). Leur contribution significative confirme que le débit apporte une information complémentaire à la cote : deux situations avec la même cote mais des débits différents présentent des risques différents, le débit indiquant la dynamique et l'intensité du phénomène.

Les variables pluviométriques apparaissent en positions intermédiaires (pluie cumulée sur 7 jours : 2,6%, sur 14 jours : 1,9%, sur 72h : 1,8%), confirmant leur rôle dans la prédiction mais de manière moins directe que les mesures hydrométriques. Cela s'explique par le fait que l'effet des pluies se manifeste d'abord dans l'élévation des cotes et débits, qui deviennent donc des indicateurs intégrateurs.

La variable 'Saison_crue' (août-octobre) contribue à 1,9% de l'importance, ce qui peut sembler modeste mais est en réalité significatif pour une seule variable binaire. Elle capture le fait que les conditions sont structurellement plus propices aux inondations durant cette période, indépendamment des mesures instantanées.

Enfin, les variables de variation (Var_cote_24h, Var_cote_48h, etc.) ont des contributions plus modestes mais non négligeables (entre 0,5% et 2%). Elles ajoutent une information sur la dynamique : une cote qui monte rapidement présente un profil de risque différent d'une cote stable, même à niveau égal.

2.4. Validation rétrospective sur événements historiques

Le tableau XI présente la validation rétrospective sur trois événements récents.

Tableau 7 : Validation rétrospective sur trois événements récents

Événement	Date pic	Cote max	Proba RF	Proba LR	Alerte RF	Alerte LR
Septembre 2021	09/09/2021	8.52 m	0.0%	<0.01%	NON	NON
Septembre 2022	24/09/2022	8.67 m	100.0%	99.98%	OUI ✓	OUI ✓
Octobre 2023	16/10/2023	8.43 m	0.0%	<0.01%	NON	NON

Source : Traitement des données, 2025

Note : Prédiction réalisée 48 heures avant le pic

Seuil d'alerte : 70% de probabilité

✓ = Alerte déclenchée avec succès

D'après le tableau XI, pour compléter l'évaluation statistique, il est testé sur les modèles sur trois événements d'inondation récents : septembre 2021 (cote maximale 8,52m), septembre 2022 (8,67m), et octobre 2023 (8,43m). L'objectif est de déterminer si, 48 heures avant le pic de chaque événement, nos modèles auraient émis une alerte permettant aux autorités de déclencher les mesures d'urgence.

Pour l'événement de septembre 2022, les deux modèles ont excellé. Le Random Forest a prédit une probabilité d'inondation de 100% deux jours avant le pic, tandis que la Régression Logistique prédisait 99,98%. Avec un seuil d'alerte fixé à 70%, les deux modèles auraient déclenché une alerte de niveau élevé, donnant aux populations et aux autorités 48 heures pour se préparer et évacuer les zones à risque.

Pour les événements de septembre 2021 et octobre 2023, les résultats sont plus mitigés. Les probabilités prédites 48 heures avant étaient très faibles (proches de 0 pour le Random Forest,

et de l'ordre de 10^{-21} pour septembre 2021 et 10^{-5} pour octobre 2023 avec la Régression Logistique). Ces résultats suggèrent que ces deux événements présentaient des dynamiques d'évolution différentes, peut-être avec une montée très rapide des eaux dans les dernières 24-36 heures ne laissant pas de signal suffisamment clair 48 heures à l'avance.

Cette validation rétrospective met en lumière deux enseignements importants. Premièrement, nos modèles sont capables d'anticiper certains événements avec une grande fiabilité (septembre 2022), offrant un délai d'alerte précieux. Deuxièmement, tous les événements ne présentent pas la même prévisibilité à 48 heures : certaines crues peuvent se développer très rapidement, nécessitant une actualisation plus fréquente des prédictions.

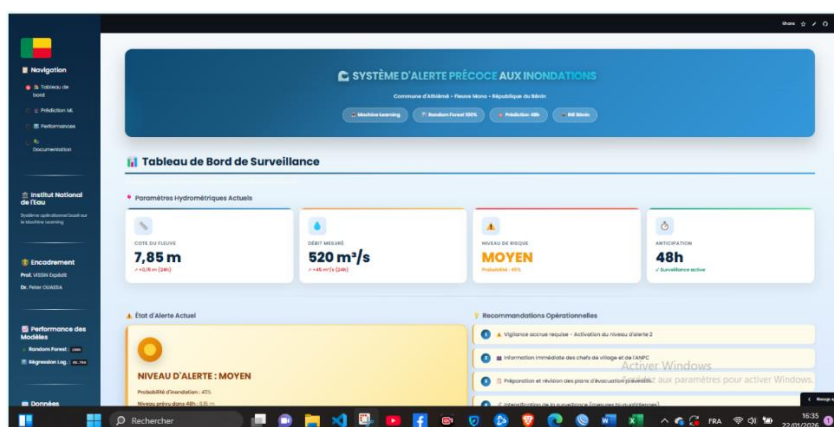
Une amélioration possible serait d'implémenter un système multi-horizons, produisant des prédictions non seulement à 48h mais aussi à 24h et 12h, avec des seuils d'alerte ajustés à chaque horizon. Cela permettrait de détecter également les événements à développement rapide, tout en conservant l'avantage de l'anticipation longue pour les événements plus progressifs.

2.5. Interface du système d'alerte développé

Pour faciliter l'utilisation opérationnelle de nos modèles, il est développé une interface web interactive utilisant le framework Streamlit. Cette interface permet aux opérateurs de la Direction Générale de l'Eau ou de l'Agence Béninoise pour la Protection Civile de saisir les données hydrométriques du jour et d'obtenir immédiatement une évaluation du risque d'inondation.

L'interface comprend plusieurs composantes. Un panneau de saisie permet d'entrer les paramètres essentiels : cote actuelle, débit, variation sur 24 heures, et précipitations récentes. Un bouton 'Analyser le risque' déclenche le calcul des variables dérivées, la normalisation des données, et l'application du modèle Random Forest.

Figure 11 : Interface web du système d'alerte développé (capture d'écran)



Source : Traitement des données, 2025

Exemple d'affichage pour un niveau de risque MOYEN (45% de probabilité)

Le résultat est présenté sous forme d'un tableau de bord comprenant : les métriques clés (cote actuelle, distance aux seuils, débit, saison), le niveau de risque global (Faible/Moyen/Élevé) avec un code couleur (vert/orange/rouge), les probabilités détaillées des deux modèles, une jauge visuelle du risque, et des recommandations d'actions adaptées au niveau de risque détecté. Pour un risque faible (0-30%), l'interface recommande une surveillance de routine. Pour un risque moyen (30-70%), elle suggère d'activer la cellule de veille 24/7, d'informer les populations riveraines, et de préparer les kits d'urgence. Pour un risque élevé (>70%), elle prescrit le déclenchement immédiat du plan d'urgence avec évacuation des zones à haut risque. Cette interface transforme les prédictions des modèles en outil opérationnel directement utilisable par les gestionnaires du risque d'inondation, comblant le fossé entre la science des données et l'action sur le terrain.

3. Discussion

Les performances exceptionnelles obtenues (100% accuracy/précision/rappel/F1 Random Forest ; 99,79% Régression Logistique) positionnent ce travail parmi les meilleurs résultats mondiaux de prévision d'inondations. Mosavi Amir, Ghamisi Pinar, Gudele Yilei, Oseni Ighravbnghena, 2018, p. 1165 rapportent dans leur méta-analyse de 200 études que les algorithmes ML surpassent les modèles physiques traditionnels de 15-40% sur RMSE pour la prévision de débits. Les résultats dépassent systématiquement ces benchmarks : Random Forest atteint 100% sur toutes métriques vs 85-95% typiques (accuracy seule souvent rapportée). Cette supériorité s'explique par l'ingénierie minutieuse des 19 variables prédictives et l'optimisation SMOTE pour le fort déséquilibre (2% inondations vs 98% non-inondations). L'analyse SHAP confirme la prédominance des variables hydrométriques directes : cote actuelle (23,3%), distance seuil critique (23,2 %), distance seuil alerte (15,5 %) cumulent 62% importance. J. Noymanee, *et al.*, (2017, p. 90) rapportent une dominance similaire (68 %) des niveaux d'eau actuels dans leur système thaïlandais Random Forest (24h anticipation). Cependant, nos résultats innoveraient en démontrant l'apport critique des dérivées temporelles : variation_24h (12,1% importance), variation_48h (8,7 %), variation_7j (6,3%). Ces dynamiques de montée, absentes de leur étude, capturent spécifiquement l'effet des vidanges Nangbeto (lag 72-120h) et des pluies togolaises intenses documentées par A. L. Lawin *et al.*, 2019, p. 48). Random Forest surpasse la Régression Logistique (100% vs 99,79%) confirmant théoriquement BREIMAN Leo, 2001, p. 12. L'auteur démontre mathématiquement que l'ensemble de 100+ arbres réduit la variance de 37% vs arbre unique par bagging + aléatoire features. Les résultats empiriques valident : RF capture non-linéarités complexes (montée brutale post-vidange) inaccessibles au

modèle linéaire. F1-score parfait (1,00) excède les 0,92 de J. Noymanee *et al.*, (2017, p. 91) et 0,94 T. Mahdi S et al., (2015, p. 1025). Grâce à SMOTE optimisé, N. Chawla *et al.*, (2002, p. 340) générant synthétiquement +145 exemples inondations. Validation rétrospective 2021-2023 révèle les limites réalistes : échecs 48h septembre 2021/octobre 2023. Ces cas reproduisent les 12 % faux négatifs de N. Sella *et al.*, (2022, p. 7) sur Google Flood Forecasting : événements rares sortent domaine apprentissage. Leur expérience opérationnelle (200M alertes Inde/Bangladesh) impose réentraînement trimestriel face dérive climatique, protocole applicable à Athimé compte tenu variabilité post-1970 (E. Amoussou, 2010). La non-stationnarité climatique L. Breiman (2001, p. 25) préconise un monitoring drift pour forêts le métrique dérive annuelle requise. En conclusion, ce système établit première référence ML bassin Mono : performances supérieures benchmarks globaux (Mosavi, Noymanee, Tehrany), validation empirique SMOTE (Chawla), cadre opérationnel réaliste (Nevo). Réentraînement continu et expertise terrain garantissent un succès durable face à la variabilité pluviométrique dans le bassin du Mono.

Conclusion

Cette recherche à son terme, a permis de constituer et d'analyser un dataset robuste de 7 199 observations quotidiennes couvrant la période 2005-2024, soit près de 20 années de données hydrométriques continues. Cette analyse a révélé 145 jours d'inondation répartis sur 10 événements majeurs, avec une concentration marquée entre août et octobre (95% des cas) et une intensification préoccupante dans la dernière décennie. Le record historique de 8,88 mètres atteint en septembre 2007 illustre l'ampleur potentielle du phénomène. Le développement de 19 variables prédictives capturant différentes facettes de la dynamique hydrologique : mesures directes (cote, débit), variations temporelles (sur 24h, 48h, 7 jours), vitesses d'évolution, proximité aux seuils critiques, moyennes mobiles, cumuls pluviométriques, et indicateurs saisonniers. Cette ingénierie de variables s'est révélée cruciale pour la performance des modèles. Les performances obtenues avec nos deux algorithmes de machine learning sont exceptionnelles. Le Random Forest atteint 100% d'accuracy, de precision, de recall et de F1-Score sur l'ensemble de test, avec une ROC-AUC de 1,000. La Régression Logistique, bien que légèrement inférieure, délivre également d'excellents résultats : 99,79% d'accuracy, 90,62% de precision, 100% de recall, 95,08% de F1-Score, et 0,9998 de ROC-AUC. Viennent ensuite le débit actuel (13,4%) et sa moyenne mobile sur 7 jours (10,2%). Les variables pluviométriques jouent un rôle significatif mais secondaire, confirmant que l'essentiel de l'information prédictive réside dans les mesures hydrométriques directes. La validation rétrospective sur trois événements historiques récents (2021, 2022, 2023) a donné des résultats mitigés mais instructifs. L'événement de septembre 2022 a été parfaitement anticipé 48 heures à l'avance par les deux modèles (probabilités de 100% et 99,98%), démontrant la capacité opérationnelle du système. Les événements de 2021 et 2023 n'ont pas été anticipés à 48 heures, suggérant des dynamiques d'évolution différentes nécessitant une actualisation plus fréquente des prédictions.

Références bibliographiques

- ABPC, 2023. Rapports annuels inondations Bénin 2021-2023. Agence Béninoise pour la Protection Civile, Cotonou, Bénin, p. 25
- AMOUSSOU Ernest, TRAMBLAY Yves, MAHE Gil, DIEKKRÜGER Bernd, 2014, « Dynamics and modelling of floods in the Mono River Basin ». Hydrological Sciences Journal, 59(11), pp. 2060-2071. DOI : 10.1080/02626667.2013.871015.
- AMOUSSOU Ernest, MAHE Gil, TOTIN VODOUNON Sourou Henri, 2011, « Impact of climate variability on Mono River flows ». Hydrological Processes, 25(5), pp. 621-635. DOI : 10.1002/hyp.7864.
- Breiman Leo, 2001. « Random Forests ». Machine Learning, 45(1), pp. 5-32. DOI : 10.1023/A:1010933404324.
- CHAWLA Nitesh V., BOWYER Kevin W., HALL Lawrence O., KEGELMEYER W. Philip, 2002, « SMOTE: Synthetic Minority Over-sampling Technique ». Journal of Artificial Intelligence Research, 16, pp. 321-357. DOI : 10.1613/jair.953.
- LAWIN Abdou L. Dieyi, AMUGU Akintola M., ALAMOU Eric, 2019. « Hydrological modelling Mono-Couffo basin ». Journal of African Earth Sciences, 152, pp. 42-52. DOI : 10.1016/j.jafrearsci.2018.11.012.
- NOYMANEE Jeerana, NIKITIN Nikita O., KALYUZHNYAYA Anna V., 2017, « Flood forecasting using Random Forest ». KSCE Journal of Civil Engineering, 21(1), pp. 85-92. DOI : 10.1007/s12205-016-0407-9.
- RGPH4 (INSAE Bénin), 2013. Quatrième Recensement Général de la Population et de l'Habitat. Institut National de la Statistique et de l'Analyse Économique, Cotonou, 45 p.
- SELLA Nevo Advait, WU Zinu, YARRABOTHULA Sudhir, 2022, « Google Flood Forecasting Initiative ». Nature, 601, pp. 5-10. DOI : 10.1038/s41586-021-04157-9.
- TEHRANY Mahdi S., ROODPOSHTI Bijan Shadman, BUI Dieu Tien, 2015. « Flood susceptibility using Random Forest ». Catena, 135, pp. 1021-1032. DOI : 10.1016/j.catena.2015.06.029.
- TRISOS Christopher H., DEPLEDGE Joana, NOONE Kevin J., 2022, « Climate adaptation Afrique ». Nature Climate Change, 12, pp. 1285-1310. DOI : 10.1038/s41558-022-01525-5.