

Utilisation de l'IA pour l'aide au *Mapping* de Paramètres de Vol dans le Cadre de la Gestion de Données Aéronautiques chez Airbus Helicopters

Thomas Calvet¹, Caroline Del-Cistia-Gallimard¹, Antoine Delaunay¹,
Adil Soubki², Pierre-Loïc Maisonneuve¹, Ammar Mechouche¹

¹ Airbus Helicopters, Aéroport International Marseille-Provence, 13700 Marignane, France
{antoine.a.delaunay, caroline.c.del-cistia-gallimard, thomas.calvet,
pierre-loic.maisonneuve, ammar.mechouche}@airbus.com

² Airbus Operations SAS, 316 Route de Bayonne, CS 83172, 31027 Toulouse, France
adil.soubki@airbus.com

Résumé

Cet article présente une approche innovante basée sur l'IA pour résoudre le défi de mapping des paramètres de vol chez Airbus Helicopters. La qualité des analyses big-data est fortement conditionnée par les processus d'ingénierie des données. En effet, les paramètres enregistrés en vol (séries temporelles multivariées) peuvent différer d'une version logicielle embarquée à une autre et d'un type d'hélicoptère à un autre. Le but du travail présenté ici est d'apporter un support aux data ingénieurs par la génération automatique de paramètres génériques uniques et faciles à identifier.

Mots-clés

Apprentissage machine, mapping de paramètres de vol, ingénierie de la donnée, support utilisateurs, application industrielle, hélicoptères.

Abstract

This paper presents an innovative machine learning approach to address the challenge of flight parameter mapping at Airbus Helicopters. The manual mapping of flight parameters—which can differ in name or unit across embedded software versions and helicopter types—is a critical yet error-prone, time-consuming, and tedious industrial bottleneck that hinders large-scale data analysis. Our work aims to support data engineers and improve the reliability of this process by leveraging both parameter content (time series) and their documentation using AI.

Keywords

Machine learning, flight parameter mapping, data engineering, user support, industrial application, helicopters.

1 Introduction

L'exploitation de données d'usage opérationnel des produits est un enjeu majeur pour les entreprises de conception et de production industrielles. Depuis près d'une ving-

taine d'années, Airbus Helicopters équipe ses produits de systèmes d'acquisitions de données embarquées, appelés Health and Usage Monitoring System (HUMS), et collecte massivement des données de vol. Plus particulièrement, des milliers de paramètres d'état de la machine et de contexte d'opération sont enregistrés en continu par le Flight Data Continuous Recorder (FDCR). Les séries temporelles qui en résultent sont traitées au sol et exposées aux différents métiers internes au travers de solutions d'analyse big data [2] et d'outils d'investigations adaptés aux experts aéronautiques [3, 11]. Aujourd'hui, les informations relayées par l'analyse des données issues du HUMS sont stratégiques pour l'entreprise : investigations d'événements pour le support au client, statistique d'usage et retours d'expériences pour adapter le design d'hélicoptères, amélioration de la maintenance et maintenance prédictive.

La qualité des analyses et la valorisation des données sont fortement conditionnées par les processus d'ingénierie et de gestion des données. L'implémentation de processus d'amélioration de la qualité des données dans les entreprises est donc un enjeu majeur [12]. Chez Airbus Helicopters, les données collectées évoluent au fil des boucles de maturation des systèmes d'acquisition embarqués et des différents hélicoptères : changement de nommage, d'unité, suppression ou ajout de paramètres, évolution des capteurs, etc. L'un des processus d'ingénierie des données clés pour les analyses consiste à harmoniser et standardiser les données en associant aux paramètres enregistrés un paramètre générique dont la définition est invariante entre les différents types d'hélicoptères et leurs versions embarquées, comme illustré sur la Figure 1. Sans cette étape, les analyses de données massives se retrouvent restreintes à une version d'acquisition, limitant la capacité statistique nécessaire à une prise de décision de confiance. A ce jour, cette étape repose sur une analyse croisée des spécifications des systèmes d'acquisition et des définitions de paramètres physiques exploitables par les analystes. Chronophage et sujet à des erreurs, ce processus ne répond pas aux enjeux indus-

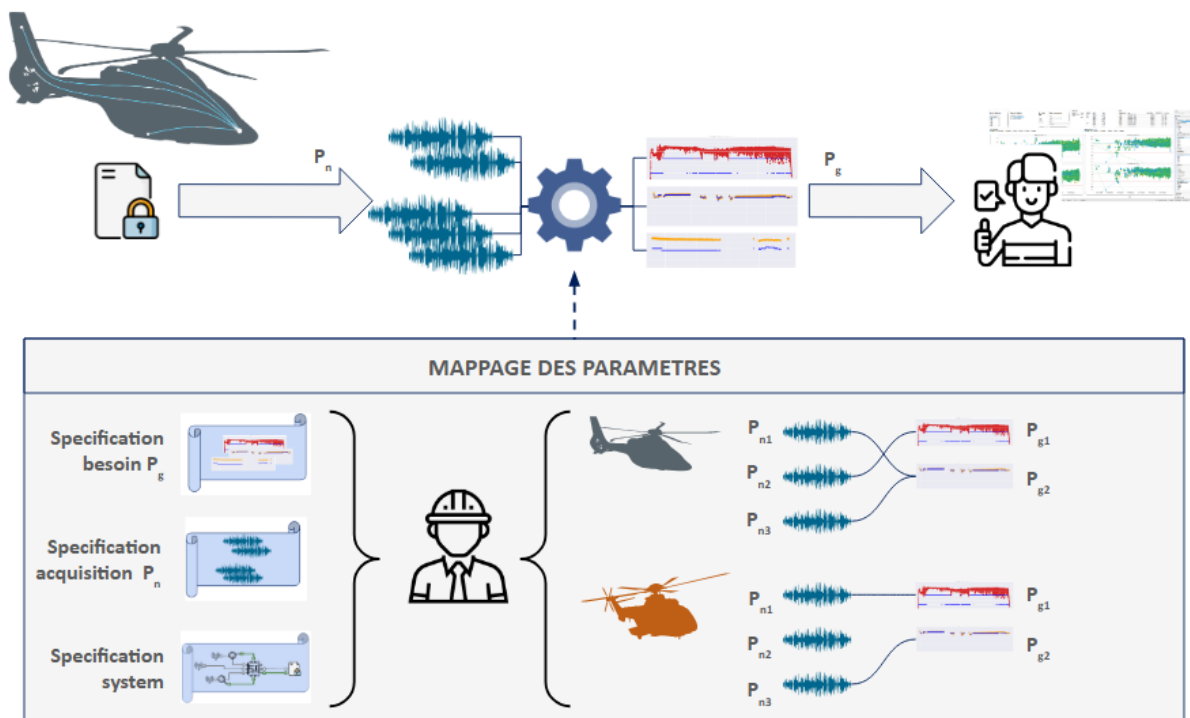


FIGURE 1 – Illustration du processus de *mapping* des paramètres.

triels.

L'apprentissage machine et l'intelligence artificielle ouvrent la voie à des avancées méthodologiques majeures. D'un côté, la compréhension des caractéristiques fondamentales des séries temporelles en permet la caractérisation précise dans un espace à dimension réduite [6, 10, 7]. De l'autre, le traitement du langage naturel permet la fouille documentaire pour contextualiser les données. Dans cet article nous proposons une approche innovante combinant ces deux aspects :

- un *mapping* comportemental par analyse de similarité entre les propriétés des séries temporelles brutes acquises sur un vol et les propriétés des paramètres génériques sur une base de référence,
- un *mapping* sémantique par la comparaison des spécifications du système d'acquisition utilisé pour l'enregistrement, notamment la nomenclature technique (souvent non-intelligible), avec un référentiel de *mapping* sur les versions historiques.

Nous démontrons la pertinence d'une telle approche et discutons l'intégration de cette méthode dans le processus de *mapping* existant. L'effort engagé pour permettre aux ingénieurs de comprendre les résultats et valider le *mapping* est mis en avant. Enfin, nous évoquons les problèmes rencontrés et les opportunités ouvertes par cette approche.

2 Contexte Industriel

Les données HUMS répondent à deux enjeux stratégiques qui imposent deux contraintes sur le système d'information. D'un côté, à chaque fin de vol les données doivent être ex-

ploitable en quelques heures afin d'informer les services de support au client d'éventuels événements à traiter. De l'autre côté, les données reçues de chaque hélicoptère alimentent les plateformes big data pour des analyses sur l'ensemble de l'historique. Des processus d'harmonisation des données sont donc essentiels mais ne doivent pas perturber les chaînes de traitement automatiques des données.

Le déploiement d'un système d'acquisition pour une version sur un hélicoptère spécifique est tracé et associé à une spécification des données enregistrées. Cette traçabilité assure la connaissance des paramètres natifs (P_n), à savoir les séries temporelles brutes, enregistrés pour chaque vol. En revanche, les modèles d'acquisition et la gamme de produits d'Airbus Helicopters ont largement évolué depuis les premiers systèmes HUMS développés dans les années 1990. Le concept de paramètre générique (P_g), à savoir un nom de paramètre harmonisé sur l'ensemble de la gamme, est introduit pour permettre d'étudier l'évolution physique de manière cohérente sur l'historique des données. Le *mapping* des paramètres consiste à associer chaque (groupe de) paramètre(s) natif(s) à un paramètre générique.

Pour résoudre l'enjeu d'uniformisation des séries temporelles, le *mapping* des paramètres dans le système d'information est configurable pour chaque vol. Le référentiel de *mapping* des paramètres est tenu à jour par des ingénieurs Airbus Helicopters. L'utilisation d'applications dédiées au *mapping* simplifie le travail de configuration et permet d'identifier les changements. Néanmoins, le travail de résolution des conflits reste manuel. Les ingénieurs des données doivent éprouver les spécifications, interviewer les métiers et définir en conséquence le *mapping* des

paramètres.

3 Méthodologie Proposée

Nous proposons une méthode hybride, combinant l'analyse du comportement intrinsèque des signaux des paramètres natifs et l'analyse sémantique de leur nomenclature. Cette méthode reproduit et systématise le raisonnement d'un expert cherchant à résoudre le *mapping* d'une nouvelle version du système d'acquisition. Au-delà de la haute précision, cette solution est robuste, scalable et produit des résultats interprétables, permettant de garantir la confiance des utilisateurs.

3.1 Représentation Hybride des Paramètres

Le cœur de notre approche consiste à transformer chaque paramètre P en une représentation vectorielle hybride de taille fixe, concaténant deux vecteurs de caractéristiques. Un vecteur comportemental V_c , représentant les propriétés statistiques et dynamiques extraites de la série temporelle $S(P)$ du paramètre, et un vecteur sémantique V_s extrait de la nomenclature $N(P)$ du paramètre en utilisant des techniques de traitement du langage naturel comme Term Frequency-Inverse Document Frequency (TF-IDF), adaptées au vocabulaire technique.

3.1.1 Axe Comportemental : Signature Statistique du Signal

L'axe comportemental permet de réduire la dimension du problème en synthétisant les séries temporelles $S(P)$, longues et de durées variables, par un vecteur synthétique et informatif V_c . Pour construire V_c , une analyse univariée est d'abord réalisée permettant de retenir vingt-huit composantes couvrant les différentes facettes d'une série temporelle, comme ses propriétés statistiques (moyenne, médiane, écart-type, asymétrie, etc.), sa dynamique temporelle (autocorrélation, pente) et sa complexité (entropie).

Enfin, les données FDCR ont des distributions très hétérogènes sur des ordres de grandeurs différents, par exemple la vitesse de rotation du rotor principal est mesurée en pourcentage du maximum spécifié, alors que l'altitude est mesurée en pieds.

3.1.2 Axe Sémantique : Exploitation de la Nomenclature

Le second volet de la modélisation hybride repose sur l'exploitation systématique des définitions issues de l'ingénierie système. Chaque paramètre acquis en vol a fait l'objet d'un processus formel de spécification allant de la définition générique de l'observable (objet, unité, fréquence, etc) jusqu'à son implémentation dans les logiciels embarqués. Bien que les spécifications de conception soient exhaustives, elles sont particulièrement complexes à exploiter : formats documentaires, bases de données propres aux outils de spécifications, langage naturel, etc. Pour cette première implémentation, nous nous sommes donc basé sur le critère

le plus directement exploitable et discriminant : le nom du paramètre $N(P)$.

L'objectif de la représentation sémantique est de modéliser cette information textuelle sous la forme d'un vecteur de caractéristiques numériques V_s . Le traitement automatique de ce champ se fait en trois étapes.

- **Tokenisation métier** : segmentation de la nomenclature en composants logiques, appelés tokens, correspondant aux concepts d'ingénierie sous-jacents.
- **Vectorisation TF-IDF** : attribution d'une pondération statistique à chaque token pour isoler les termes les plus significatifs au sein du corpus global.
- **Intégration d'un dictionnaire métier** : utilisation d'un dictionnaire spécifique pour traiter les synonymes et acronymes propres au domaine aéronautique pour capturer les similarités sémantiques qui échappent à une analyse purement lexicale.

Cette méthodologie permet de quantifier la proximité sémantique entre paramètres, complétant l'analyse comportementale par une structure ancrée dans la logique de conception du système.

3.2 Modélisation par Apprentissage Supervisé

L'hybridation des deux axes d'analyse est réalisée au sein d'un modèle d'apprentissage supervisé, capable d'apprendre à pondérer intelligemment les informations comportementales et sémantiques pour définir le *mapping*. Plus spécifiquement, le modèle déployé effectue une tâche de classification supervisée en différenciant les couples de paramètres (P_n, P_g) correspondants, classés comme 1 (positifs), des autres couples de paramètres, classés comme 0 (négatifs). L'architecture de la solution dans son ensemble est représentée sur la Figure 2.

3.2.1 Préparation des Données

Plusieurs centaines de milliers d'heures de vol ré-échantillonnées à la demi-seconde sont disponibles sur le lac des données d'Airbus Helicopters. Cette base big data est utilisée pour générer les vecteurs V_c . Un tel vecteur est généré pour chaque paramètre pour chaque vol. Cent à trois cent mille points (selon le type d'hélicoptère) de données sont ainsi obtenus. Pour chaque paramètre, autant de vecteurs V_s sont construits qu'il y a de version de système embarqués. Le jeu de données final est constitué du couple hybride (V_s, V_c) . Le partitionnement des données en 70 % entraînement, 10 % validation et 20 % test, garantissant a minima plusieurs dizaines de milliers de points pour chaque étape, suit une contrainte de robustesse stricte : une version logicielle donnée (et donc une représentation V_n) est exclusivement assignée à un seul sous-ensemble. Cette stratégie garantit l'évaluation de la capacité du modèle à généraliser ses prédictions sur des configurations avioniques totalement inconnues.

La performance d'un modèle supervisé dépend de manière critique de la qualité de son jeu de données d'entraînement.

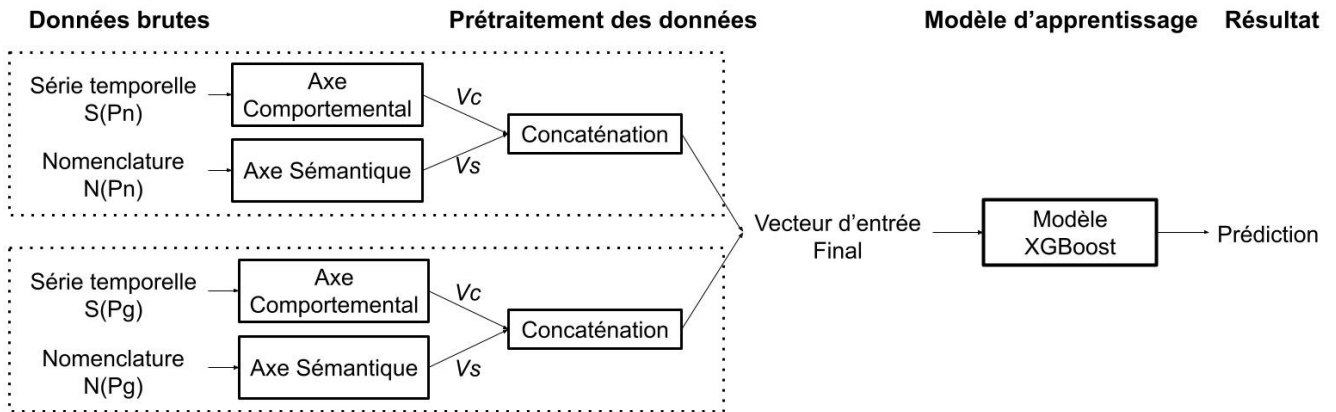


FIGURE 2 – Illustration de l'architecture globale de la solution proposée.

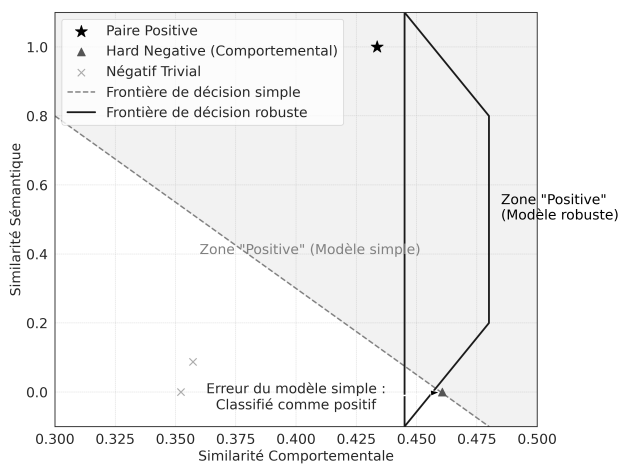


FIGURE 3 – Identification de l'espace des "hard negatives" dans le plan de similarité comportementale (axe des x) et similarité sémantique (axe des y). Le détour noir montre l'espace choisi.

Dans le cas du *mapping* des paramètres, pour une paire cible (P_n, p_g) il existe des milliers de paires fausses (P_n', P_g) . Une fraction importante de ces paires fausses constitue des exemples triviaux (par exemple la paire température extérieure - booléen de logique vol/sol). Leur maintien dans l'entraînement abrutit le modèle et dégrade les performances. Pour surmonter cet écueil, les classes sont pondérées et une stratégie agressive d'échantillonnage de "hard negatives" est mise en place. Elle consiste à inclure intentionnellement dans le jeu de données des paires négatives "difficiles" : des paramètres qui ne correspondent pas mais qui présentent de fortes similarités, que ce soit sur le plan comportemental (par exemple deux températures de circuits distincts) ou sémantique (par exemple des noms très proches). La Figure 3 montre le domaine du plan des similarités conservé pour construire les jeux d'entraînement. En forçant le modèle à se confronter à ces cas pièges, ce dernier construit des frontières de décision plus fines et plus robustes, améliorant significativement sa

capacité de généralisation.

3.2.2 Choix du Modèle, Réglage des Hyperparamètres et Stratégie d'Entraînement

Les modèles ensemblistes basés sur des arbres de décision sont reconnus pour leur robustesse et leur performance sur des données tabulaires hétérogènes, ainsi que leur facilité de compréhension et d'analyse. Un modèle *Extreme Gradient Boosting* (XGBoost) est donc testé dans cette analyse. Le modèle XGBoost [4], opérant sur un principe d'agrégation séquentielle ("boosting") où des arbres de décision sont construits itérativement pour corriger les erreurs des précédents, est retenu comme principal candidat pour l'architecture de la solution. Le choix de XGBoost est justifié par plusieurs avantages techniques alignés avec les contraintes du projet :

- **Capacité d'optimisation avancée** : intégration de mécanismes de régularisation L1/L2 contrôlant la complexité du modèle et réduisant le surapprentissage.
- **Gestion efficace des variables corrélées** : gestion robuste de la présence de caractéristiques corrélées, fréquent dans notre cas avec des features comportementales comme la moyenne et la médiane d'un même signal.
- **Performance empirique supérieure** : Démonstration de performances supérieures dans de nombreux benchmarks de classification sur des données tabulaires.
- **Explicabilité intégrée** : Facilitation de l'utilisation d'outils d'explicabilité (comme SHAP [9] ou l'importance des caractéristiques), étant un atout majeur dans un contexte industriel où chaque décision de l'algorithme doit pouvoir être justifiée auprès des experts métier.

Les bonnes pratiques de construction de modèle d'apprentissage sont appliquées en suivant une procédure rigoureuse et reproductible, conçue pour garantir la robustesse des résultats. La recherche des hyperparamètres (nombre d'arbres, profondeur maximale, taux d'apprentissage, etc)

optimaux des modèles est automatisée via l'utilisation du framework Optuna [1] s'appuyant sur des stratégies bayésiennes. Cette approche est couplée à une validation croisée afin d'éviter le surajustement et d'assurer la capacité de généralisation du modèle. En complément, une technique d'arrêt anticipé, interrompant l'entraînement si les performances sur l'ensemble de validation n'évoluent plus, est employée pour prévenir le surapprentissage.

Par ailleurs, la construction d'un grand nombre de paires négatives, comme décrit dans la section 3.2.1, entraîne un déséquilibre entre le nombre de paires positives et de paires négatives. Afin de remédier à ce déséquilibre, pouvant impacter les performances du modèle sur la classe minoritaire, des techniques d'équilibrage sont appliquées, comme la pondération des classes lors du calcul de la fonction de coût du modèle d'apprentissage.

4 Expérimentations et Résultats

L'architecture hybride proposée est appliquée sur des cas d'usage réels afin de quantifier ses performances, d'analyser sa robustesse et de décrire son intégration opérationnelle.

4.1 Protocole Expérimental

La solution de *mapping* vise à traiter les nouvelles versions d'avionique. Dans une démarche de validation rigoureuse, une attention forte est portée sur la capacité du modèle à généraliser sur des données nouvelles.

Métriques d'Évaluation. Les performances du modèle de classification binaire ont été quantifiées à l'aide de quatre métriques standards.

- *Précision* : proportion de correspondances correctes (vrais positifs) parmi les positifs classés par le modèle.
- *Rappel* : proportion de correspondances possibles identifiées par le modèle.
- *F1-score* : moyenne harmonique de la précision et du rappel.
- *Aire sous la Courbe ROC (AUC)* : capacité globale du modèle à discriminer les paires positives des paires négatives.

Analyse de Sensibilité. Afin d'optimiser le comportement du modèle, l'ajout d'un paramètre de mémoire système a été analysé. Nous avons observé que l'ajout du paramètre *is_historical_match*, attestant de l'existence de la paire (P_n, P_g) analysée dans le référentiel des *mappings*, augmente le risque de surapprentissage et réduit la capacité du modèle à généraliser en mettant un accent trop fort sur l'historique.

Garantie de la Reproductibilité. Enfin, un effort particulier a été consacré pour garantir la reproductibilité de l'ensemble de la chaîne expérimentale. Plusieurs mesures

ont été mises en place pour assurer la traçabilité et l'auditabilité des résultats. Celles-ci incluent une journalisation systématique des configurations d'exécution (plages de sessions, listes de paramètres, versions de code), un versioning de tous les artefacts de données intermédiaires et finaux, ainsi qu'une conception modulaire d'un pipeline de traitement qui sépare clairement chaque étape pour faciliter les ré-exécutions ciblées. Ces dispositions assurent que les résultats présentés sont vérifiables et peuvent être reproduits de manière fiable.

4.2 Analyse Quantitative des Performances

L'évaluation du modèle sur des jeux de données indépendants confirme la validité de l'architecture proposée.

Performance Globale. Le tableau 1 résume les performances obtenues et démontre la haute efficacité de l'approche proposée. Le système atteint une exactitude de **93.5 %** et un ROC AUC de **0.948**, témoignant d'une haute capacité de discrimination entre paires de paramètres. Avec une précision de **84.2 %** et un rappel de **72.7 %**, le modèle automatise l'identification de près des trois quarts des correspondances existantes avec un haut niveau de fiabilité. Le F1-Score de 0.78 qui en résulte confirme le bon équilibre entre ces deux derniers indicateurs. Ces résultats confirment que le système peut automatiser une part très significative du processus de *mapping* avec un haut niveau de fiabilité, validant ainsi l'apport majeur de l'approche par gradient boosting.

TABLE 1 – Mesure des métriques de performance du modèle XGBoost.

Exactitude	Précision	Rappel	F1-Score	ROC AUC
0.9348	0.8421	0.7273	0.7805	0.9483

Analyse des Limites et des Cas d'Erreur. L'analyse des résidus indique que les erreurs de classification se concentrent sur les paramètres à faible variance, tels que les alarmes binaires dont l'occurrence est rare ou nulle dans les données opérationnelles. Dans ces configurations, la nomenclature ne permet pas l'identification et l'absence de signature comportementale stable limite la capacité de détection de l'algorithme. Alors qu'un volume significatif de données est mis à disposition, la représentativité des paramètres dans les données opérationnelles reste un enjeu. Plus particulièrement lorsque le paramètre caractérise des événements redoutés que le design industriel limite ou élimine.

4.3 Déploiement et Validation Opérationnelle

Au-delà des métriques quantitatives, la valeur de la solution réside dans son intégration effective dans les processus

métier. Le modèle est un assistant à l'expert métier.

4.3.1 Idée Générale

Le Paradigme "Human-in-the-Loop". L'outil a été conçu selon une logique de "human-in-the-loop" (l'humain dans la boucle). L'expert métier conserve le contrôle final : il peut valider, corriger ou rejeter les suggestions de l'algorithme. Cette interaction est fondamentale, car elle permet non seulement de traiter les cas les plus ambigus où le modèle est incertain, mais aussi de collecter un retour précieux. Chaque décision de l'expert est réinjectée dans le système pour l'entraînement futur, créant ainsi une boucle d'amélioration continue qui garantit l'évolution et l'affinement constant du modèle.

Garantie de Confiance et Adoption. Cette approche collaborative est un facteur clé pour assurer la confiance et l'adoption de l'outil. En positionnant le modèle comme un assistant intelligent qui automatise les tâches répétitives et met en lumière les cas complexes, plutôt que comme une "boîte noire" qui impose ses décisions, le système s'intègre naturellement dans les flux de travail existants. Le déploiement de cet outil constitue ainsi un résultat à part entière, marquant le passage d'un projet de recherche à une solution opérationnelle et évolutive.

4.3.2 Intégration dans une Interface Web Homme-Machine

Afin de faciliter son usage opérationnel, le système de *mapping* a été intégré dans une interface web dédiée, qui constitue le point de contact principal entre le modèle et les ingénieurs HUMS. Cette interface ne se contente pas d'afficher les résultats ; elle présente les correspondances suggérées, les scores de confiance associés, ainsi que des visualisations et des indicateurs statistiques pour appuyer la décision de l'expert.

La page principale de l'interface web présentant les résultats du *mapping* est visualisée sur la Figure 4.

Dans cet exemple on peut voir que le pourcentage de paramètres directement mappés en comparant avec l'historique (ce qui est considéré aujourd'hui), les *mappings* suggérés permettant d'élargir la sélection des paramètres considérés pour la nouvelle version, les cas dupliqués pour lesquels deux propositions de *mapping* sont disponibles. Enfin les cas pour lesquels le *mapping* n'est pas possible, c'est-à-dire que la base de données actuelle ne contient pas d'informations pour le *mapping* de ce types de paramètres, sont aussi mis en avant.

Chaque page de l'interface web se concentre sur un type de résultat du *mapping*. Par exemple, la page "Conflicts to resolved" (Figure 5) répertorie les *mappings* pour lesquels les axes sémantique ("historical") et comportemental ("behavioral") entrent en contradictions. On peut noter que pour l'axe comportemental, un indice de confiance est systématiquement associé au résultat. Dans le premier exemple, la divergence entre les deux prédictions ne porte que sur le numéro de l'équipement (1 ou 2). Cette subtilité peut in-

duire une certaine ambiguïté lors de l'analyse comportementale, ce qui explique néanmoins l'indice de confiance élevé de 91%. À l'inverse, dans le second exemple, le score de confiance de 12% sur l'axe comportemental reflète fidèlement l'incertitude du modèle et corrobore la prédiction erronée. Il est toutefois intéressant de souligner que, malgré cette faible confiance, le paramètre identifié reste proche de la valeur attendue.

L'approche choisie s'inscrit dans une logique de "human-in-the-loop" : l'expert peut confirmer, corriger ou rejeter une suggestion. Ces retours alimentent ensuite le système afin de renforcer ses performances futures, créant une boucle d'amélioration continue. Cette intégration pratique illustre la volonté de développer non pas un outil théorique isolé, mais un système réellement utilisable, évolutif et aligné avec les pratiques opérationnelles d'Airbus Helicopters.

5 Discussion

Au-delà des performances quantitatives présentées au chapitre précédent, une discussion des choix méthodologiques et de leurs implications est nécessaire pour apprécier la portée de ce travail. Cette section a pour but de prendre du recul sur les résultats, en explicitant les compromis réalisés lors de la conception du système et en positionnant sa contribution dans un contexte plus large d'ingénierie des données.

5.1 Compromis entre Performance et Explicabilité

Le choix méthodologique a privilégié l'explicabilité et la confiance du modèle face à la recherche de performance absolue. L'axe comportemental repose sur un ensemble restreint de vingt huit comportements statistiques experts, sélectionnés pour leur signification physique explicite (moyenne, dispersion, entropie). Bien que des approches "deep learning" ou des solutions packagées du marché telles que Rocket/MiniRocket [5] pourraient offrir des gains en performance, leur nature "boîte noire" compliquerait la justification des résultats auprès des experts métiers. Les progrès notables en transparence des résultats d'IAs, allant jusqu'à l'analyse des chaînes de pensées des LLMs [8], restent spécifiques au domaine de la donnée et n'atteignent pas encore les métiers industriels. Ces approches sont à suivre pour les améliorations du *mapping* de paramètres.

Dans un contexte industriel critique comme l'aéronautique, un modèle légèrement moins performant mais dont les décisions sont crédibles, vérifiables et justifiables est préférable à un modèle opaque. Le système proposé incarne cette philosophie : chaque correspondance suggérée peut être expliquée par des propriétés tangibles du signal, assurant ainsi la confiance nécessaire au déploiement opérationnel.

5.2 Approche Univariée : Robustesse contre Exhaustivité

Le second arbitrage méthodologique majeur a été le choix d'une décomposition univariée des séries temporelles pour alimenter une classification multivariée. Il est important de reconnaître la limite inhérente à ce choix : il ignore les corrélations potentiellement riches d'information qui existent entre les signaux (par exemple, entre le booléen d'allumage de la climatisation pour différencier température intérieure et extérieure). Cette approche entraîne une perte notable d'information et donc de performance.

Cependant, cette approche apporte un gain substantiel en robustesse, un critère prioritaire dans notre contexte. Les avantages opérationnels sont décisifs :

- **Adaptabilité aux évolutions logicielles** : Le modèle reste applicable quelle que soit la configuration des capteurs, qui peut évoluer d'une version logicielle à l'autre (capteurs ajoutés, modifiés ou supprimés).
- **Robustesse face aux données manquantes** : Contrairement à un modèle multivarié qui exigerait la présence de tous les signaux pertinents, notre système reste fonctionnel même si certains capteurs sont temporairement inactifs ou en panne.
- **Simplicité de validation** : Une décision basée sur les caractéristiques d'un seul signal est beaucoup plus simple à vérifier et à valider pour un expert qu'une décision issue d'interactions complexes entre de multiples paramètres.

Dans un contexte d'innovation, où pour la première fois la décision métier de *mapping* sera appuyée par une décision machine, la fiabilité face aux cas opérationnels connus a retenu notre attention. L'intégration d'algorithmes multivariés dès la synthèse des paramètres pour améliorer les performances est à confronter à une expérience opération-

nelle de la méthode proposée.

5.3 Implications pour l'Ingénierie des Données

Enfin, il est essentiel de souligner que la contribution de ce travail dépasse la simple conception d'un modèle de classification. La problématique de départ n'était pas seulement technique, mais également organisationnelle et stratégique. Le processus de *mapping* manuel constituait un verrou technologique majeur et un goulot d'étranglement fondamental, empêchant l'exploitation agile et à grande échelle des données de vol.

En proposant une solution d'automatisation robuste, scalable et intégrée dans une boucle "*human-in-the-loop*", ce projet contribue à la résolution d'un problème fondamental d'ingénierie des données. Il ne fournit pas seulement des prédictions, mais il fiabilise et accélère une étape critique de la chaîne de valeur analytique. Ce faisant, il constitue un socle essentiel à la stratégie big data de l'entreprise, en garantissant la qualité et la cohérence des données nécessaires à toutes les analyses de flotte, qu'il s'agisse de maintenance prédictive, de suivi de la sécurité ou d'aide à la conception.

6 Conclusion

Cet article présente une approche robuste et explicable pour automatiser le processus de *mapping* entre paramètres natifs et paramètres génériques afin d'harmoniser les données de vol des hélicoptères de la flotte Airbus Helicopters. La méthode repose sur une synthèse comportementale univariée des signaux et sémantique de la nomenclature des paramètres acquis en vol, et la combinaison des informations obtenues par apprentissage machine multivarié. L'éva-

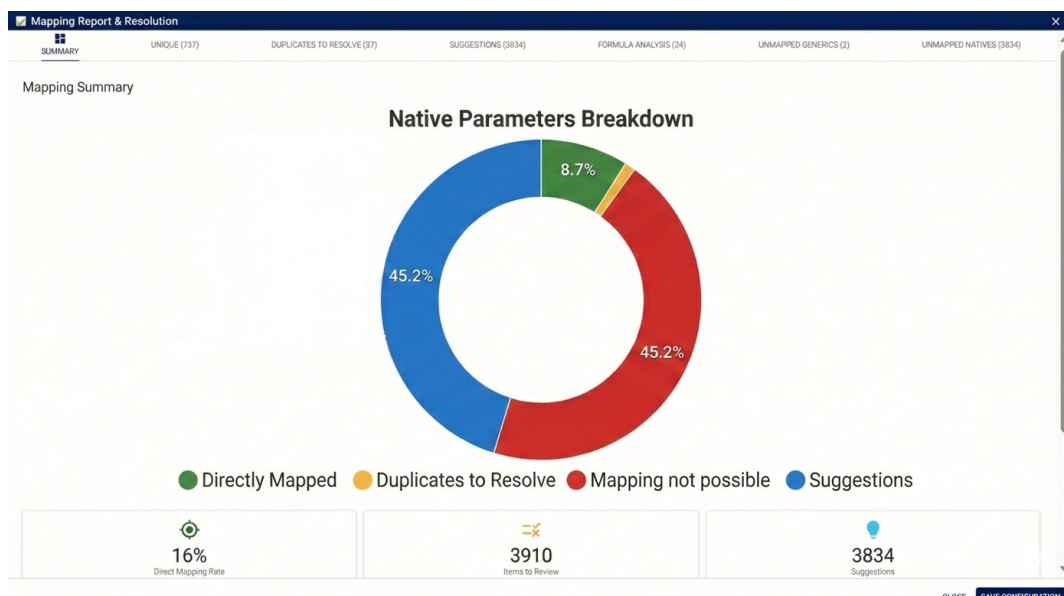


FIGURE 4 – Tableau de bord de l'interface web présentant les résultats du *mapping*.

Behavioral Mapping Report					
SUMMARY	ALIGNED PREDICTIONS (44)	CONFLICTS TO RESOLVE (78)	NEW SUGGESTIONS (332)	HISTORICAL ONLY (160)	ANALYSIS FAILED (2553)
Conflicts to Resolve (Historical & Behavioral predictions differ)					
Native Parameter ↑	Behavioral Prediction ↑	Confidence ↑	Historical Mapping ↑	Resolve Conflict	
DC1GeneratorVoltageValue_VMS_C	ELECTRICAL_DC2_GENERATOR_VOLTAGE_VALUE	90.87%	ELECTRICAL_DC1_GENERATOR_VOLTAGE_VALUE	KEEP BEHAVIORAL	KEEP HISTORICAL
DC2GeneratorCurrentValue_VMS_C	ELECTRICAL_DC2_GENERATOR_VOLTAGE_VALUE	11.55%	ELECTRICAL_DC2_GENERATOR_CURRENT	KEEP BEHAVIORAL	KEEP HISTORICAL

FIGURE 5 – Exemple de conflit entre les résultats donnés par l’axe sémantique et l’axe comportemental.

luation du système, exécuté via un pipeline de traitement distribué, a montré des résultats encourageants, avec une précision moyenne de 84 %.

L’intégration d’une boucle “human-in-the-loop” et le déploiement d’une interface web ont permis de replacer l’outil dans un cadre opérationnel concret. Ces avancées contribuent directement à la réduction de l’effort manuel permettant l’intégration plus rapide d’une nouvelle version de système d’acquisition embarqués, une meilleure traçabilité des décisions et une réduction du risque d’erreur. Néanmoins, certaines limites demeurent. Les performances du modèle sont intrinsèquement faibles pour les paramètres variant peu, ou qui visent à détecter les situations que le design des produits tend à limiter. Les paramètres complexes construits par des formules de paramètres natifs ne sont pas gérés.

Ces constats ouvrent la voie à plusieurs perspectives, que ce soit sur la méthodologie utilisée pour le *mapping* ou son intégration. Tout d’abord, la gestion des paramètres à faible variance pourrait être améliorée via l’utilisation de méthodes d’augmentation des données. Des approches multivariées pour la synthèse des signaux pourraient également être explorées afin de traiter les cas les plus complexes. D’autre part, la lisibilité des résultats et la capacité des ingénieurs à les interpréter est rendue possible par l’utilisation de concepts intermédiaires formels, robustes et connus. L’intégration d’outils d’apprentissage plus profond pour gérer plus en amont les corrélations ou augmenter la capacité décisionnelle du modèle est conditionnée par l’adhérence des métiers industriels aux nouvelles méthodes d’explicabilité des IAs.

Références

[1] Takuya Akiba, Shotaro Sano, Toshihiko Yanase, Takeru Ohta, and Masanori Koyama. Optuna : A next-generation hyperparameter optimization framework.

In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, KDD ’19, page 2623–2631, New York, NY, USA, 2019. Association for Computing Machinery.

- [2] Mechouche Ammar, Daouayry Nassia, and Camerini Valerio. Helicopter big data processing and predictive analytics : feedback perspectives. *ERF2019 0035*, 2019.
- [3] Nicolas Brisset. Big data tooling and usage perspectives in the airbus helicopters test center. In *European Test and Telemetry Conference (ettc 2022)*, pages 175–181, 01 2022.
- [4] Tianqi Chen and Carlos Guestrin. Xgboost : A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 785–794, 08 2016.
- [5] Angus Dempster, François Petitjean, and Geoffrey I Webb. ROCKET : Exceptionally fast and accurate time series classification using random convolutional kernels. *Data Mining and Knowledge Discovery*, 34(5) :1454–1495, 2020.
- [6] Johann Faouzi. *Time Series Classification : A Review of Algorithms and Implementations*, chapter 2. IntechOpen, London, 2024.
- [7] Hassan Ismail Fawaz, Germain Forestier, Jonathan Weber, Lhassane Idoumghar, and Pierre-Alain Muller. Deep learning for time series classification : a review. *Data Mining and Knowledge Discovery*, 33 :917 – 963, 2018.
- [8] Tomek Korbak, Mikita Balesni, Elizabeth Barnes, Yoshua Bengio, Joe Benton, Joseph Bloom, Mark Chen, Alan Cooney, Allan Dafoe, Anca Dragan, Scott Emmons, Owain Evans, David Farhi, Ryan Greenblatt, Dan Hendrycks, Marius Hobbhahn, Evan Hubinger,

Geoffrey Irving, Erik Jenner, Daniel Kokotajlo, Victoria Krakovna, Shane Legg, David Lindner, David Luan, Aleksander Mądry, Julian Michael, Neel Nanda, Dave Orr, Jakub Pachocki, Ethan Perez, Mary Phuong, Fabien Roger, Joshua Saxe, Buck Shlegeris, Martín Soto, Eric Steinberger, Jasmine Wang, Wojciech Zaremba, Bowen Baker, Rohin Shah, and Vlad Mikulik. Chain of thought monitorability : A new and fragile opportunity for ai safety, 2025.

- [9] Scott Lundberg and Su-In Lee. A unified approach to interpreting model predictions, 2017.
- [10] Jorge Rocha, Sandra Oliveira, and Cláudia Viana. *Time Series Analysis - Recent Advances, New Perspectives and Applications*. IntechOpen, 05 2024.
- [11] Malet Vincent and Soubki Adil. Ease datalake fishing with language models. In *European Test and Telemetry Conference (ettc 2025)*, 2025.
- [12] Agung Wahyudi, George Kuk, and Marijn Janssen. A process pattern model for tackling and improving big data quality. *Information Systems Frontiers*, 20 :457–469, 2018.