

Proposition d'un sujet de thèse en cotutelle 2026/2027

A. Porteur du sujet à l'Université Libanaise

1. Nom : JABER
2. Prénom : Ali
3. Titre (Prof, HDR,) : Prof
4. Laboratoire : L'ARICoD
5. Adresse Web :
6. Etablissement : Université Libanaise
7. Adresse Web : <https://www.ul.edu.lb/>
8. Domaines d'expertise :
 - Réseaux de Capteurs sans Fils, Intelligence Artificielle, Sciences de l'Information, Qualité de Service, Analyse de Données Volumineuses
9. Publications importantes en relation avec le sujet proposé :
 - Fadlallah, H., Kilany, R., Haber, M., & **Jaber, A.** (2023, December). CTXDQ: An automated context-driven data quality assessment. In 2023 IEEE 4th International Multidisciplinary Conference on Engineering Technology (IMCET) (pp. 32-37). IEEE.
 - Al Tfaily, F., Hazimeh, H., El Takach, A., El Zein, H., Harb, H., & **Jaber, A.** (2025, October). SAGE: Semantic Adaptation of Graph Embeddings for Arabic Entity Linking. In 2025 IEEE/ACS 22nd International Conference on Computer Systems and Applications (AICCSA) (pp. 1-5). IEEE.
 - Zaraket, K., Harb, H., Bennis, I., **Jaber, A.**, & Abouaissa, A. (2024). Hyper-Flophet: A neural Prophet-based model for traffic flow forecasting in transportation systems. *Simulation Modelling Practice and Theory*, 134, 102954.
 - **Jaber, A.**, & Guarnieri, F. (1998). Un système d'agents logiciels pour favoriser la coopération entre des systèmes d'aides à la décision dédiés à la gestion de crise. *Systèmes multi-agents, de l'interaction à la socialité: Actes de JFIADSMA'98 Sixièmes Journées Francophones d'Intelligence Artificielle et Systèmes Multi-Agents*, ISBN-9782746227941.
10. Adresse Web de votre page personnelle :
11. Adresse e-mail : ali.jaber@ul.edu.lb

B. Partenaire à l'ULCO :

1. a. Nom : BOUNEFFA
b. Nom : AHMAD
2. a. Prénom : Mourad
b. Prénom : Adeel
3. Titre (Prof, HDR, ...) :
 - a. Professeur des Universités - Informatique
 - b. Maître de Conférences - Informatique
4. Laboratoire : Laboratoire d'Informatique Signal et Image de la Côte d'Opale (LISIC)
5. Etablissement : Université du Littoral Côte d'Opale, France
6. Adresse Web : <https://lisic-prod.univ-littoral.fr/>
7. Domaines d'expertise :
 - eXplicabilité de l'Intelligence Artificielle,
 - Ingénierie de Connaissances,
 - Systèmes de recommandations,
 - Raisonnement Sémantique,
 - Probabilistic Graphical models,
 - Knowledge graphs,
 - Rule-based Constraints.
8. Publications importantes en relation avec le sujet proposé :
 - Mohammad Choib, Moncef Garouani, **Mourad Mohamed Bouneffa, Adeel Ahmad**. IoT-AID: Leveraging XAI for Conversational Recommendations in Cyber-Physical Systems. 27th International Conference on Enterprise Information Systems, INSTICC : Institute for Systems and Technologies of Information, Control and Communication, Apr 2025, Porto, Portugal. pp.671-679
 - Moncef Garouani, **Mourad Mohamed Bouneffa, Adeel Ahmad**, Mohamed Hamlich. Version [2 . 0] - [AMLBIID: An auto-explained Automated Machine Learning tool for Big Industrial Data]. SoftwareX, 2023, 23, pp.101444.
 - Moncef Garouani, **Adeel Ahmad, Mourad Mohamed Bouneffa**. Explaining Meta-Features Importance in Meta-Learning Through Shapley Values. 25th International Conference on Enterprise Information Systems (ICEIS 2023), Apr 2023, Prague, France. pp.591–598.

9. Adresse Web de votre page personnelle : <https://dr-adeel.github.io/>
10. Adresse e-mail :
 - a. mourad.bouneffa@univ-littoral.fr
 - b. adeel.ahmad@univ-littoral.fr

C. Description du sujet de thèse proposé : (3 à 5 pages)

1. Discipline : Informatique
2. Titre : IA explicable grâce à la modélisation de l'espace d'état fonctionnel et à la rétropropagation sémantique : vers une explicabilité centrée sur l'humain dans les systèmes multi-agents
3. Résumé :

L'intégration croissante de l'intelligence artificielle (IA) dans la surveillance de la santé mentale via des dispositifs portables présente un défi majeur : garantir la transparence, l'interprétabilité et la confiance dans les évaluations pilotées par l'IA. Cette recherche propose de relever ce défi en développant un cadre de modélisation de l'espace d'état fonctionnel (FSS) pour représenter les indicateurs de santé mentale - tels que la variabilité de la fréquence cardiaque, la conductance de la peau, les cycles de sommeil ou encore les mouvements - dans un modèle cognitif interprétable. En étendant la rétropropagation aux espaces sémantiques, cette étude permet aux modèles d'IA d'expliquer la façon dont les signaux physiologiques sont corrélés avec les conditions de santé mentale, rendant ainsi leurs processus de prise de décision plus transparents et cliniquement significatifs.

Par ailleurs, cette recherche explore l'intelligence collective décentralisée dans le contexte de l'IA portable, en permettant aux modèles IA multi-capteurs de collaborer efficacement à partir de sources de données variées (EEG, montres intelligentes, trackers de fitness, etc.). Le cadre proposé sera validé à travers des applications concrètes, notamment pour la détection en temps réel du stress et de l'anxiété. En alliant modélisation cognitive, raisonnement sémantique et coopération décentralisée, cette étude contribue à renforcer la confiance des utilisateurs et des cliniciens dans l'IA appliquée à la santé mentale.

4. Sujet :

a. Description du sujet (contexte scientifique, description du problème, Objectifs,) :

La complexité croissante des modèles d'intelligence artificielle (IA) [1][2] représente un défi majeur pour garantir la transparence et l'explicabilité, en particulier dans des domaines de prise de décision critique tels que la santé, la finance, le droit et les systèmes autonomes. Cette recherche vise à relever ce défi en développant un cadre de Modélisation de l'Espace d'État Fonctionnel (FSS) [3], une approche formalisée de la modélisation de la cognition de l'IA à travers des distances sémantiques mesurables et des transitions d'état structurées.

La modélisation de l'espace d'état fonctionnel (FSS), également appelée modélisation fonctionnelle centrée sur l'humain (HCFM), représente les domaines cognitifs et physiques sous forme de réseaux réversibles basés sur des métriques. Elle définit les états fonctionnels d'un système comme des nœuds dans un graphe, avec des transitions entre ces états représentées par des arêtes. Cette approche vise à créer une « représentation sémantique complète » du domaine d'un système, permettant une communication et une collaboration claires. La modélisation FSS cherche à cartographier un « espace de définition du problème » et un « espace de solution », nécessitant des efforts d'alignement complexes. Elle repose sur la modélisation sémantique et des solutions d'IA basées sur des graphes métriques pour quantifier la similarité entre les états fonctionnels. En représentant les processus de prise de décision de l'IA au sein d'un espace conceptuel structuré, la modélisation FSS offre une méthode réversible et basée sur des métriques pour comprendre comment les systèmes d'IA raisonnent, offrant ainsi un moyen de démystifier les modèles "boîte noire".

b. Approche méthodologique :

Les modèles traditionnels d'IA fonctionnent souvent de manière isolée, tandis que les problèmes complexes du monde réel nécessitent des interactions collaboratives et multi-agents. Cette recherche examinera comment les structures de hypergraphes peuvent faciliter une prise de décision plus interprétable et efficace de l'IA en permettant des échanges sémantiques plus riches entre les agents d'IA. En établissant des chemins de raisonnement structurés au sein des

hypergraphes, les systèmes d'IA peuvent participer à un partage des connaissances plus transparent et responsable, réduisant ainsi les risques de prise de décision opaque et inexpliquée.

Un élément central de cette recherche est l'extension de la rétropropagation vers les espaces sémantiques, en introduisant la Rétropropagation Sémantique [4] comme un mécanisme visant à améliorer l'alignement et l'interprétabilité de l'IA. L'apprentissage traditionnel basé sur les gradients repose sur des techniques d'optimisation numériques qui manquent de raisonnement sémantique explicite. En incorporant des structures sémantiques dans le processus d'apprentissage, cette approche améliore la capacité de l'IA à généraliser des concepts d'une manière à la fois compréhensible pour les humains et fonctionnellement cohérente. Cette avancée est particulièrement pertinente pour garantir que les modèles d'IA utilisés dans les applications réelles [5] maintiennent la cohérence logique et la transparence dans leur prise de décision.

c. Résultats attendus :

Pour garantir l'applicabilité pratique de ces avancées théoriques, les modèles proposés seront validés empiriquement dans des domaines réels, en particulier dans les systèmes de santé. En appliquant la modélisation FSS, la rétropropagation sémantique et l'intelligence basée sur des hypergraphes dans ces domaines, cette recherche vise à démontrer des améliorations mesurables de la transparence de l'IA, de la confiance des utilisateurs et de la supervision humaine. Ces contributions fourniront un cadre fondamental pour le développement de systèmes d'IA plus interprétables, responsables et éthiquement alignés, garantissant que l'IA puisse être intégrée de manière sûre et efficace dans des applications sociétales critiques.

Le travail de recherche proposé pourrait répondre à des questions telles que : “Les systèmes d'IA basés sur des hypergraphes de plus haut niveau peuvent-ils faciliter une fusion multi-capteurs interprétable pour des analyses personnalisées de la santé mentale ? Comment ces méthodes d'IA explicables peuvent-elles être appliquées pour détecter les signes précoces d'anxiété, de dépression et de stress en temps réel, garantissant un déploiement éthique et centré sur l'utilisateur ?”

Cette recherche se concentre principalement sur l'amélioration de l'interprétabilité de la surveillance de la santé mentale pilotée par l'IA grâce au développement d'un cadre FSS qui

mappe les indicateurs physiologiques et comportementaux - tels que la variabilité de la fréquence cardiaque, la conductance de la peau, les patterns de sommeil et les mouvements - dans un espace cognitif structuré. En étendant la rétropropagation aux espaces sémantiques, l'étude vise à permettre aux modèles d'IA de fournir des explications claires sur la manière dont les signaux physiologiques sont corrélés avec les conditions de santé mentale, améliorant ainsi la transparence et la confiance. De plus, la recherche explore l'intelligence collective décentralisée dans l'IA portable, garantissant une collaboration fluide entre les modèles multi-capteurs d'IA à travers diverses sources de données, telles que l'EEG, les montres intelligentes et les trackers de fitness. Afin de valider l'efficacité de ces méthodologies, des applications réelles seront examinées, en particulier dans des contextes cliniques et grand public, démontrant l'impact de l'IA explicable sur l'amélioration de la confiance des cliniciens et des utilisateurs dans la surveillance de la santé mentale assistée par IA. Une application clé de cette recherche est la détection en temps réel du stress et de l'anxiété à l'aide de l'IA portable, fournissant aux individus des informations interprétables et exploitables sur leur bien-être mental.

Cette thèse relie les avancées de pointe dans la modélisation de la cognition de l'IA aux défis urgents de l'explicabilité de l'IA dans le monde réel. En combinant rigueur mathématique, modélisation cognitive et raisonnement décentralisé de l'IA, cette recherche ouvrira la voie à des systèmes d'IA non seulement plus puissants mais aussi plus transparents, alignés sur l'humain et éthiquement robustes.

d. Calendrier :

Cette recherche vise à développer des méthodes d'IA explicables pour la surveillance de la santé mentale à travers des dispositifs portables. Elle propose un cadre de modélisation basé sur l'Espace d'État Fonctionnel (FSS) pour représenter les indicateurs physiologiques et comportementaux dans un espace cognitif interprétable, tout en intégrant des méthodes avancées de rétropropagation sémantique. Le projet inclut également une exploration de l'intelligence collective décentralisée pour améliorer la collaboration entre les différents modèles d'IA et dispositifs de collecte de données.

Année 1 : Modélisation Cognitive et Développement du Cadre FSS

1. Phase de Recherche Initiale et de Revue de Littérature (Semestre 1)

- Examiner l'état de l'art en matière d'IA explicable, de santé mentale, et de modèles de détection des signes physiologiques et comportementaux.
- Développement du Cadre Fonctionnel de l'Espace d'État (FSS)
- Modéliser les indicateurs physiologiques et comportementaux de la santé mentale (comme la variabilité de la fréquence cardiaque, la conductance de la peau, les modèles de sommeil et les mouvements) dans un espace cognitif structuré à l'aide de FSS.
- Élaborer un cadre de mesure des distances sémantiques et des transitions d'état dans le modèle FSS.

2. Extension de la Rétropropagation vers des Espaces Sémantiques (Semestre 2)

- Adapter la rétropropagation pour qu'elle fonctionne dans des espaces sémantiques, permettant aux modèles d'IA d'expliquer la manière dont les signaux physiologiques sont corrélés aux conditions de santé mentale.
- Assurer l'intégration des connaissances expertes dans le modèle hybride connexionniste/symbolique.
- Conception et Intégration des Dispositifs Portables pour la Collecte des Données
- Sélectionner et configurer les dispositifs portables tels que les montres intelligentes, les trackers de fitness et les EEG pour collecter des données pertinentes sur les indicateurs de santé mentale.
- Tester et intégrer les capteurs physiques et comportementaux.

Année 2 : Intégration et Validation des Modèles AI Explicables

1. Développement de l'Intelligence Collective Décentralisée (3-4 mois)

- Rechercher et mettre en œuvre des structures de graphes hypergraphiques pour permettre l'échange sémantique entre les agents IA multi-capteurs.

- Assurer la collaboration transparente entre différents dispositifs de collecte de données et modèles d'IA.
- 2. Application à la Détection en Temps Réel de l'Anxiété et du Stress (4-5 mois)
 - Appliquer le cadre FSS et les approches de rétropropagation sémantique à la détection de l'anxiété, du stress et des symptômes de dépression.
 - Intégrer les dispositifs portables pour fournir des résultats en temps réel et interprétables aux utilisateurs.
- 3. Validation Empirique et Réalisation d'Études de Cas (4-5 mois)
 - Tester les modèles développés dans des contextes cliniques et consommateurs, en mettant l'accent sur les résultats des applications réelles (détection du stress et de l'anxiété).
 - Effectuer des analyses empiriques pour évaluer l'impact de l'IA explicable sur la confiance des cliniciens et des utilisateurs.
- 4. Analyse des Résultats et Amélioration du Modèle (2-3 mois)
 - Analyser les performances des modèles en termes de précision, d'explicabilité et de transparence.
 - Affiner les mécanismes de rétropropagation et d'intelligence collective pour améliorer l'efficacité du système d'IA dans les applications en santé mentale.

Année 3 : Consolidation des Modèles et Applications Pratiques

1. Optimisation et Généralisation des Modèles d'IA (4-5 mois)
 - Développer des outils de sélection et d'optimisation des modèles d'IA explicables dans des contextes multidimensionnels.
2. Extension des Applications à d'Autres Domaines (3-4 mois)
 - Étudier et appliquer les modèles aux domaines de la finance, du droit, ou des systèmes autonomes pour démontrer leur adaptabilité dans divers contextes.
 - Vérifier l'efficacité des systèmes dans d'autres environnements pour renforcer leur robustesse.

3. Validation Finale et Tests Longitudinaux (4-5 mois)

- Mener des tests longitudinaux pour valider les résultats sur le long terme, notamment en suivant l'évolution des symptômes chez les utilisateurs ayant utilisé des dispositifs portables.
- Examiner l'impact de l'IA explicable sur la gestion des conditions de santé mentale, en s'assurant qu'elle reste éthique et centrée sur l'utilisateur.

4. Rédaction et Défense de la Thèse (2-4 mois)

- Présenter les résultats obtenus en mettant l'accent sur l'avancement des méthodes d'IA explicables dans les applications pratiques de santé mentale.
- Finaliser la rédaction de la thèse, synthétiser les résultats et les contributions de la recherche.

e. Mots clés :

eXplainable Artificial Intelligence (XAI), Ingénierie de Connaissances, Rule-based Constraints, Healthcare, Modélisation de l'Espace d'État Fonctionnel, Rétropropagation Sémantique, raisonnement sémantique.

f. Possibilité de financement (Justificatif éventuel) :

Co-financement LISIC/ULCO 50%

g. L'état du sujet dans le laboratoire et l'équipe d'accueil :

En cohérence avec les travaux menés par l'équipe IC (Ingénierie des Connaissances) du LISIC, cette recherche s'inscrit dans la continuité des explorations menées sur l'IA explicable. Elle vise à intégrer les connaissances expertes dans un cadre hybride associant modèles connexionnistes et approches symboliques. Ce projet repose sur des travaux antérieurs, dont la thèse consacrée à un système de recommandation pour les systèmes cyber-physiques.

Cette thèse présente un framework conversationnel de recommandation pour les systèmes cyber-physiques (CPS), conçu pour assister les ingénieurs et les experts du domaine dans la configuration

complexe de ces systèmes. En s'appuyant sur une base de connaissances sémantique et le traitement du langage naturel, il permet de traduire des objectifs de haut niveau en recommandations techniques optimisées. Le système innove par l'intégration d'une boucle de clarification interactive pour lever les ambiguïtés et de techniques d'IA explicable pour justifier ses décisions. L'approche développée réduit significativement le temps de conception tout en renforçant la confiance des utilisateurs grâce à une transparence accrue.

Le présent projet vise à enrichir cette base en orientant les efforts vers la construction de modèles explicables impliquant des réseaux neuronaux, en particulier pour la détection de caractéristiques. La thèse contribuera à lever un verrou actuel : l'intégration d'explications compréhensibles dans des modèles de type réseaux neuronaux convolutifs (CNN), en prolongeant les capacités de l'outil développé vers des modèles plus complexes tout en conservant une explicabilité utile aux spécialistes du domaine.

h. Les retombées scientifiques et économiques attendues :

Cette recherche vise des retombées à la fois scientifiques, technologiques et socio-économiques. D'un point de vue scientifique, elle contribuera à faire progresser l'état de l'art en matière d'IA explicable, en particulier dans les domaines de la santé numérique et de l'intelligence multi-agent. L'introduction de concepts comme la rétropropagation sémantique ou les hypergraphes cognitifs représente un potentiel de publication élevé dans des revues internationales à comité de lecture.

Le projet permettra également de renforcer la collaboration entre le LISIC et des partenaires socio-économiques, en favorisant le transfert de technologies explicables vers les secteurs de la santé, des objets connectés, ou encore de la cybersécurité. Des partenariats industriels pourront émerger autour de la valorisation des solutions développées (logiciels, outils embarqués, protocoles de collecte de données). La possibilité de créer des modules explicatifs pour dispositifs portables connectés pourrait aussi susciter un intérêt commercial dans les filières de santé numérique ou de bien-être personnel.

Enfin, le projet contribue à la formation d'un jeune chercheur dans un domaine stratégique, en lui fournissant des compétences de pointe en IA hybride, en ingénierie des systèmes intelligents et en



Université Libanaise



analyse éthique des technologies. Ce positionnement, à la croisée des enjeux techniques et sociétaux, constitue une véritable valeur ajoutée pour l'intégration professionnelle du doctorant.