



Avis de Soutenance

THESE DE DOCTORAT

Présentée par

Madame IKRAM EL MIQDADI

Discipline : Informatique

Spécialité : Informatique

Sujet de la thèse

Détection du contenu raciste sur les réseaux sociaux

Formation Doctorale " Sciences de l'Ingénieur, Sciences Physique, Mathématiques et Informatique"

Thèse présentée et soutenue **le samedi 26 septembre 2026 à 10h** au Centre des Conférences à la Faculté des Sciences et Techniques de Fès, devant le jury composé de :

NOM ET PRÉNOM	TITRE	ÉTABLISSEMENT	
AZEDDINE ZAHI	PES	Faculté des Sciences et Techniques de Fès	Président
AICHA MAJDA	PES	Faculté des Sciences Juridiques, Economiques et Sociales de Meknès	Rapporteur
SAID BENHLIMA	PES	Faculté des Sciences de Meknès	Rapporteur
ARSALANE ZARGHILI	PES	Faculté des Sciences et Techniques de Fès	Rapporteur
KHALID ABBAD	PES	Faculté des Sciences et Techniques de Fès	Examineur
SOUFIANE HOURRI	MCH	Ecole Supérieure de Technologie de Safi	Examineur
JAMAL KHARROUBI	PES	Faculté des Sciences et Techniques de Fès	Directeur de Thèse

Laboratoire de recherche : Systèmes Intelligents et Applications

Etablissement : Faculté des Sciences et Techniques de Fès



Centre d'Etudes Doctorales : Sciences et Techniques et Sciences Médicales

Résumé de la thèse

La prolifération des discours haineux sur les plateformes numériques constitue un défi majeur pour les sociétés contemporaines. Parmi ces formes de violence en ligne, le racisme occupe une place centrale. Favorisé par l'anonymat, il se manifeste sous des formes explicites ou implicites, rendant la modération manuelle difficile et nécessitant le développement de systèmes automatiques. Cette thèse propose une étude de la détection automatique du racisme dans les réseaux sociaux dans un cadre multilingue couvrant l'anglais, le français, l'arabe standard moderne et l'arabe dialectal nord-africain. Elle s'appuie sur une approche combinant la construction de données, l'apprentissage automatique, le deep learning et l'apprentissage par transfert, afin de répondre à plusieurs défis majeurs tels que la rareté des données annotées, le déséquilibre des classes et la diversité linguistique. Dans un premier temps, une revue systématique de la littérature, menée selon le protocole PRISMA, a permis d'analyser les limites des ressources existantes et de guider la construction d'une base de données multilingue annotée à partir de Twitter et Facebook. Plusieurs modèles classiques d'apprentissage automatique, tels que SVM, KNN, Decision Trees, Random Forest et MLP, ont été étudiés, montrant que des approches simples peuvent être particulièrement efficaces dans des contextes de données limitées. Par ailleurs, une approche d'enrichissement automatique des données a été proposée, combinant Sentence-BERT et réseaux de neurones convolutifs dans un cadre de self-training, permettant d'augmenter significativement la taille du corpus tout en améliorant les performances de classification supervisée. Une comparaison entre modèles classiques et modèles profonds a également été réalisée à deux niveaux : sur la base originale et sur la base augmentée. Elle montre que SVM reste compétitif face aux approches profondes même sur des volumes de données relativement importants, tandis que CNN surpasse les autres approches lorsque le volume de données devient important. Enfin, une approche de transfer learning basée sur XLM-RoBERTa avec fine-tuning partiel a été explorée, offrant de bonnes performances pour certaines langues tout en révélant des limitations pour l'arabe. L'ensemble de ces travaux met à disposition des ressources annotées et ouvre des perspectives vers des approches plus flexibles de transfer learning et des systèmes multimodaux de détection du contenu raciste. Mots-clés : Détection automatique du racisme, discours haineux, réseaux sociaux, traitement multilingue, apprentissage automatique, deep learning, apprentissage par transfert, XLM-RoBERTa, self-training, classement de texte.