

الجمهورية الشعبية الديمقراطية الجزائرية
People's Democratic Republic of Algeria
وزارة التعليم العالي والبحث العلمي
Ministry of Higher Education and Scientific Research
المدرسة العليا للإعلام الآلي 8 ماي 1945 - سيدي بلعباس
Higher School of Computer Science
8 Mai 1945 - Sidi Bel Abbès



Thesis

To obtain the diploma of **Engineering Degree**
Field of Study: **Computer Science**
Specialization: **SIW**

Theme

AI-Based Models for Malicious URL Detection Using Metaheuristic Optimization

Presented by
Bellout Sarra

Defended on: **July 8th, 2025**
In front of the jury composed of

Dr. Anani Djihed
Dr. Khaldi Miloud
Dr. Mahammed Nadir
Dr. BabaAhmed Manel

President of the Jury
Thesis Supervisor
Co-Supervisor
Examiner

Academic Year: 2024/2025

Abstract

Phishing and other malicious URL-based attacks continue to be significant risks in cybersecurity, taking advantage of the shortcomings of conventional detection systems like blacklists and heuristic rule engines.

This study introduces an AI-driven detection framework that combines traditional machine learning algorithms (XGBoost, LightGBM) and deep learning (BiLSTM) with metaheuristic optimization methods to enhance detection precision and minimize false positives.

Our technique emphasizes the extraction and engineering of lexical and statistical information from URLs, followed by tokenization and TF-IDF vectorization. We use metaheuristics, such as Particle Swarm Optimization (PSO) and Optuna, for automatic hyperparameter optimization. This greatly enhances model efficacy and resilience. Experimental assessment of benchmark datasets verifies that the improved models attained excellent accuracy (up to 98.94%), robust generalization, and minimal false alarm rates.

The findings illustrate the capability of integrating traditional AI models with optimization techniques to develop effective and scalable phishing detection systems. The suggested methodology is efficient and transparent, providing a viable answer to dynamic cybersecurity challenges.

Keywords: Phishing Detection, Malicious URLs, Machine Learning, Deep learning, Metaheuristics, PSO, Optuna

Résumé

Les attaques de phishing et les URL malveillantes demeurent parmi les menaces les plus répandues en matière de cybersécurité, exploitant les failles des systèmes de détection classiques comme les listes de blocage et les règles heuristiques.

Cette étude présente un cadre de détection utilisant l'intelligence artificielle qui fusionne des algorithmes traditionnels d'apprentissage automatique (XGBoost, LightGBM) avec un modèle profond (BiLSTM), qui ont été optimisés à l'aide de techniques métaheuristiques.

La méthodologie utilisée repose sur l'extraction de caractéristiques lexicales et statistiques des URLs, suivie d'une vectorisation TF-IDF. L'optimisation des hyperparamètres est automatisée en utilisant PSO (Optimisation par Essaim Particulaire) et Optuna, ce qui conduit à une amélioration significative de la précision des modèles. Les tests expérimentaux ont démontré que les modèles optimisés peuvent atteindre une précision pouvant aller jusqu'à 98,94%, tout en conservant un faible taux de faux positifs.

Les résultats obtenus mettent en évidence le potentiel des approches hybrides qui combinent l'intelligence artificielle classique et l'optimisation pour développer des systèmes de détection efficaces et adaptables contre les menaces de phishing contemporaines.

Mots-clés : Détection de phishing, URLs malveillantes, Apprentissage automatique, Apprentissage profond, Métaheuristiques, PSO, Optuna

ملخص

تُعدّ هجمات التصيد الاحتيالي وروابط المواقع الخبيثة من بين أكثر التهديدات شيوعاً في مجال الأمن السيبراني، حيث تستغل نقاط الضعف في أنظمة الكشف التقليدية مثل القوائم السوداء والقواعد المعتمدة على التحليل اليدوي.

يقدم هذا العمل إطاراً ذكياً يعتمد على الذكاء الاصطناعي للكشف عن روابط URLs الضارة، من خلال دمج نماذج التعلم الآلي التقليدية (مثل XGBoost و LightGBM) مع نموذج تعلم عميق (BiLSTM) وذلك باستخدام تقنيات التحسين الميتاهيورستية لتحسين الأداء.

يرتكز النهج المتبع على استخراج خصائص إحصائية ولغوية من الروابط، تليها عملية تحويل نصية باستخدام TF-IDF. تم تحسين أداء النماذج تلقائياً من خلال تقنيات مثل PSO و Optuna، مما ساهم في تحقيق دقة عالية (تصل إلى 98.94%) مع معدل منخفض للإنذارات الكاذبة.

تُبرز النتائج فعالية الجمع بين النماذج التقليدية وتقنيات التحسين لبناء أنظمة كشف هجينة، فعالة وقابلة للتكيف وقادرة على التكيف مع التهديدات الحديثة.

الكلمات المفتاحية: كشف التصيد الإلكتروني، الروابط الضارة، التعلم الآلي، التعلم العميق، الخوارزميات فوق الاستدلالية، Optuna, PSO