



الجمهورية الجزائرية الديمقراطية الشعبية

RÉPUBLIQUE ALGÉRIENNE DÉMOCRATIQUE ET POPULAIRE

وزارة التعليم العالي والبحث العلمي

MINISTÈRE DE L'ENSEIGNEMENT SUPÉRIEUR ET DE LA RECHERCHE
SCIENTIFIQUE

المدرسة الوطنية العليا للتكنولوجيا والهندسة - عنابة

ECOLE NATIONALE SUPÉRIEURE DE TECHNOLOGIE ET D'INGÉNIERIE - ANNABA

Département Génie Industriel

MEMOIRE

En vue d'obtention du diplôme d'INGÉNIEUR D'ÉTAT

Domaine : Science et Technologie

Filière : Génie Industriel

Spécialité : Maintenance et Fiabilité des Systèmes Industriels

Présenté par

MOHAMED LAMINE BOUAZIZ

DAYA EL HAK SAADI

DEVELOPPEMENT D'UNE APPROCHE PREDICTION BASEE SUR LES SVM DE LA CHAINE LOGISTIQUE SUPPLY SIDE CAS D'ETUDE : ENTREPRISES DE PRODUCTION PHARMACEUTIQUE - DAR EL- BEIDA –ALGERIE

Encadré par

Dr. Imen DRISS

Co-Encadré par

Dr. Ismail BOUACHA

ENSTI Annaba

Membres du jury :

Dr. Nour Elislem KARABADJI

Président

ENSTI ANNABA

Dr. Saida GHERBI

Examineur

ENSTI ANNABA

Mme. Lamia AISSAOUI

Examineur

ENSTI ANNABA

Année 2024

Remerciement

Tout d'abord, nous nous remercions Dieu le tout puissant de m'avoir donné le courage et la patience nécessaires à mener ce travail à son terme. Nous tenons à remercier tout particulièrement Notre encadrante, Mme. Driss Imen, pour l'aide compétente qu'il nous a apportée, pour sa patience et son encouragement. Son œil critique nous a été très précieux pour structurer le travail et pour améliorer la qualité des différentes sections.

Nos vifs et chaleureux remerciements s'adressent également au Co-encadrent Mr Bouacha esmail, pour sa disponibilité, sa coopération et son encouragement continu, qui a fait tout son possible pour nous apporter le soutien et l'assistance nécessaire tout le long de préparation de ce mémoire de Fin d'Etudes. Un très grand remerciement et une très grande reconnaissance sont destinés à Mme ait Younes Nacima de nous avoir aidé à obtenir ce stage de fin d'études chez SAIDAL.

Nous remercions aussi l'ensemble des enseignants du département de Génie Industriel de l'ENSTI qui ont été toujours disponible et à l'écoute de tous les étudiants, en particulier les responsables de la spécialité de Génie Industriel. Que les membres de jury trouvent, ici, l'expression de nos sincères remerciements pour l'honneur qu'ils nous font en prenant le temps de lire et d'évaluer ce travail.

Pour finir, nous souhaitons remercier toute personne ayant contribué de près ou de loin à

La réalisation de ce travail

Dédicace

Je dédie ce travail avant tout

À mes très chers parents qui n'ont jamais cessé de m'encourager et me soutenir et j'espère que j'aurais un jour l'occasion de les remercier davantage, Que dieu leur procure bonne santé et longue vie.

À mes chers frères et sœurs

A toute ma famille et tous mes amis pour leurs encouragements.

A tous ceux qui ont contribué de près ou de loin pour que ce projet soit possible.

Je vous dis merci.

Résumé

Dans ce projet nous avons exploré l'utilisation de l'intelligence artificielle (IA) pour anticiper les besoins futurs en matières premières. Le défi majeur est la prédiction précise de la demande en matières premières dans les années à venir. Nous passons en revue les méthodes d'IA, notamment les machines à vecteurs de support (SVM), qui se sont avérées efficaces dans divers domaines. Ce travail détaille la création de notre modèle SVM pour la prédiction de quantités, avec une évaluation des performances et une comparaison avec d'autres approches. En conclusion, nous soulignons l'importance de l'IA dans la gestion des matières premières.

Mot clés : l'intelligence artificielle, machines à vecteurs de support, prédiction.

Abstract

In this project, we explored the use of artificial intelligence (AI) to anticipate future demand for raw materials. The major challenge is to accurately predict the demand for raw materials in the coming years. We review AI methods, including support vector machines (SVMs), which have proven effective in various fields. This work details the creation of our SVM model for quantity prediction, with a performance evaluation and comparison with other approaches. In conclusion, we highlight the importance of AI in commodity management.

Key words: artificial intelligence, support vector machines, prediction.

ملخص

في هذا مشروع ، استكشفنا استخدام الذكاء الاصطناعي لتوقع الطلب المستقبلي على المواد الخام. يتمثل التحدي الرئيسي في التنبؤ بدقة بالطلب على المواد الخام في السنوات القادمة. نستعرض أساليب الذكاء الاصطناعي، ولا سيما آلات ناقلات الدعم التي أثبتت فعاليتها في مختلف المجالات. وتوضح هذه الأطروحة بالتفصيل إنشاء نموذج الخاص بنا للتنبؤ بالكمية، مع تقييم الأداء والمقارنة مع الأساليب الأخرى. وفي الختام، نسلط الضوء على أهمية الذكاء الاصطناعي في إدارة السلع. الكلمات المفتاحية: الذكاء الاصطناعي، آلات ناقلات الدعم، التنبؤ.

Liste des abréviations

SVM : les machines à vecteurs de support.

RNN : réseaux neuronaux.

PSO : Particle Swarm Optimization (Optimisation par essaims de particules).

IPSO : Invasive Plant Optimization (Optimisation des plantes envahissantes).

LS-SVM : machine à vecteur de support des moindres carrés.

IA : L'intelligence artificielle.

ML : machine Learning (Apprentissage automatique).

DL : deep Learning (apprentissage profond).

SVR : Support Vector Regression.

MSE : Erreur quadratique moyenne (Root Mean Squared Error).

RBF : Fonction de base radiale (Radial Basis Function).

List des figures et tableaux

Figure 1 : location	3
Figure 2 : les départements de L'unité de Dar El Beida	4
Figure 3 : les domaine d'application de l'apprentissage automatique	7
Figure 4: l'intelligence artificielle et ses sous demain.....	11
Figure 5: la relation entre Apprentissage automatique et Apprentissage profond	12
Figure 6: code de l'importation des bibliothèques nécessaire	13
Figure 7: code de chargement de fichier dataset newdaya2017.cvs	14
Figure 8: code de l'encodage des variables	14
Figure 9 : code de préparation des donnés	15
Figure 10 : code de normalisation des données	15
Figure 11 : diviser les données	15
Figure 12 : code d'entraiment des données	16
Figure 13 : code d'évaluation des données	16
Figure 14 : code d'évaluation croisée	16
Figure 15 : code de convertir en pourcentages.....	17
Figure 16 : code de Calcule de la moyenne et l'écart type	17
Figure 17 : code de sauvegarder de model	18
Figure 18 : Organisation du contenu	20
Figure 19 : exécuté le script Python	21
Figure 20 : interface model	21
Figure 21 : Représentation de l'évolution des déférant matière entre 2023 et la réduction en 2024	26
Figure 22: l'évaluation de la performance du modèle SVM	27

Liste des tableaux

Table 1: Matériel utilisé pour création de model svm.....	13
Table 2 : Résultat de prédiction sur 2023.....	21
Table 3 : Résultat de prédiction sur 2023 Après Avoir Travaillé Avec Le Noyau RBF.....	23
Table 4 : Résultat de prédiction sur 2024.....	24
Table 5 : Table de comparaison des techniques en fonction de l'accuracy	28

Table de matière

Remerciement	I
Dédicace	II
Résumé	III
Liste des abréviations	IV
List des figures et tableaux.....	V
Table de matière	VI
Introduction générale.....	1
Chapitre I : Présentation de l'entreprise Saidal	2
Introduction	3
I.1 L'historique [1]	3
I.2 Unité de Dar El Beida.....	3
I.2.1 Historique et activité de l'unité [1]	3
I.2.2 Les principales caractéristiques de l'unité : [1].....	4
I.2.3 Organigramme du L'unité de Dar El Beida.....	4
I.3 Motivation	4
I.4 Problématique.....	5
Chapitre II : État de l'Art	6
Introduction	7
II.1 Synthèse bibliographique	7
Chapitre 03 : création d'un modèle de machine learning	10
Introduction	11
III.1 Définition de l'intelligence artificielle	11
III.2 L'intelligence Artificiel est ses sous domaines	11
III.3 Apprentissage automatique (Machine Learning)	11
III.4 Apprentissage profond (Deep Learning)	12
III.5 Les Types d'apprentissage automatique.....	12
III.6 Les techniques du d'apprentissage automatique utilisé dans ce chapitre :.....	12
III.6.1 Support Vector Machine (SVM).....	12
III.6.2 SVR (Support Vector Regression)	12
III.6.2.1 Objectif de SVR.....	13
III.6.2.2 Fonctionnement de SVR.....	13

III.7	Matériel utilisé.....	13
III.8	Les Etapes de création d'un modèle de machine learning utilisant le (SVM) et (SVR)..	13
III.8.1	Importation des Bibliothèques	13
III.8.2	Charger un fichier CSV dans un DataFrame Pandas	14
III.8.3	Encoder des variables catégorielles en valeurs numériques l'aide de LabelEncoder 14	
III.8.4	Définir les features (caractéristiques) et la target (cible) dans un DataFrame Pandas 14	
III.8.5	Normaliser les caractéristiques d'un DataFrame à l'aide de StandardScaler.....	15
III.8.6	Diviser des tableaux ou les matrices	15
III.8.7	Évaluer un modèle de régression SVM.....	16
III.8.8	Validation croisée sur un modèle SVM (Support Vector Machine)	16
III.8.9	Convertit les scores de validation croisée en pourcentages	17
III.8.10	Calcule la moyenne et l'écart type des erreurs en pourcentage obtenues précédemment	17
III.8.11	Sauvegarder plusieurs objets liés à un modèle d'apprentissage automatique.....	18
	Conclusion	18
	Chapitre IV : interprétation et discussion des résultats	19
	Introduction	20
IV.1	Présentation d'interface	20
IV.2	Fonctionnement de l'interface	20
IV.3	Discussion Et Comparaison Des Résultat de prédiction	21
IV.4	La présentation des résultats du model sur 2024.....	24
IV.5	Représentation graphique des résultats.....	26
IV.6	Représentation graphique de l'évaluation de la performance du modèle SVM.....	27
IV.7	Discussion des choix effectués	27
	Conclusion	28
	Conclusion générale	29
	Bibliographie	31

Introduction générale

Supply side joue un rôle crucial dans le fonctionnement des entreprises. Cependant, la gestion traditionnelle de l'approvisionnement de la matière première repose souvent sur l'expérience individuelle et l'intuition, sans recourir à des méthodes formelles. Dans ce mémoire, nous explorons une approche plus rigoureuse basée sur l'intelligence artificielle (IA) pour anticiper les besoins futurs en approvisionnement.

Notre travail se concentre sur l'entreprise « Saidal » où nous avons effectué mon stage. Nous avons identifié un défi majeur : la prédiction précise de la demande en matières premières dans les années à venir. L'objectif est de créer un modèle d'approvisionnement capable de prédire en quantité, tout en éliminant la dépendance à l'égard de l'expérience individuelle.

Dans le deuxième chapitre, nous passons en revue l'état de l'art des applications de l'IA pour la prédiction de l'approvisionnement de matières premières. Nous examinons les méthodes existantes, notamment les réseaux neuronaux et les machines à vecteurs de support (SVM). Ces derniers se sont avérés efficaces dans divers domaines, grâce à leur capacité à gérer des données complexes et à généraliser à partir d'échantillons limités.

Le troisième chapitre détaille la création de notre modèle de machine learning basé sur les SVM. Nous suivons une démarche à suivre pour construire notre modèle de prédiction de quantités. Ces étapes incluent la collecte et la préparation des données, la sélection des caractéristiques pertinentes, le choix du noyau SVM, l'entraînement du modèle et l'évaluation des performances.

Dans le dernier chapitre, nous présentons les résultats obtenus pour dix produits cible. Nous interprétons ces résultats en analysant les prédictions du modèle et en les comparant aux données réelles. Puis, nous traçons les courbes de quantité en fonction du temps (t) et ($t+1$) pour visualiser les tendances. Enfin, nous évaluons la performance de notre modèle SVM pour le problème de prédiction. Nous comparons également nos résultats avec d'autres approches.

Enfin, ce mémoire s'achève par une conclusion générale résume nos travaux et souligne l'importance de l'IA dans la gestion des matières premières.

Chapitre I : Présentation de l'entreprise Saidal

Introduction

Ce chapitre sera consacré en premier lieu à la présentation générale du groupe pharmaceutique Saïdal (son historique, son organisation), et en second lieu on va présenter l'unité de Dar el Beida (le lieu de stage) où nous effectuons notre stage de PFE.

I.1 L'historique [1]

En 1969 : la pharmacie centrale algérienne pour assurer le monopole de l'Etat sur l'importation, la fabrication et la commercialisation des médicaments à usage de l'homme.

En 1971 : institut d'el Harrach dans le cadre de la mission de production avec l'acquisition des unités de Biotic et Pharmal en 2 étapes (1971 puis 1975).

À la suite de la restructuration de la PCA, sa branche de production fut érigée en entreprise nationale de production pharmaceutique (ENPP) par décret 82/161 promulgué en avril 1982.

En 1985, l'ENPP a changé de dénomination pour devenir « Saïdal », et en 1989 et suite à la mise en œuvre des réformes économiques, Saïdal est devenue une EPE avec une autonomie de gestion a été choisie parmi les premières entreprises nationales pour acquérir le statut de société par actions.

En 1997, la société Saïdal a mis en œuvre un plan de restructuration qui s'est traduit par sa transformation en groupe industriel le 2 février 1998 auquel sont rattachées trois filiales (Pharmal, Biotic et Antibiotical) issues de cette restructuration.

I.2 Unité de Dar El Beida



Nom : Saïdal.

Fax : 023 92 01 78.

Téléphone portable : 0558 05 26 96.

Email : dg.saidalgroupe@saidalgroupe.dz.

Site Web: <http://www.saidalgroupe.dz>.

Adresse : Route de wilaya n°11 BP 141 Dar El Beida – Alger.

Figure 1 : location

I.2.1 Historique et activité de l'unité [1]

L'unité de Dar el Beida est considérée comme l'unité la plus ancienne des unités de Pharmal. Cette unité existe depuis 1958, elle appartenait au laboratoire Français LABAZ avant sa nationalisation en 1970, elle a été rattachée à 51%, et en 1976 à 100% par l'ex PCA ce qui a donné lieu aux transformations suivantes :

- Agrandissement de l'unité de 3600m² à 6600 m²
- La création des produits pharmaceutiques en Algérie.

- Extension du magasin de stockage
- Modernisation des chaînes et des ateliers

I.2.2 Les principales caractéristiques de l'unité : [1]

- ✓ Une longue période d'expérience dans le domaine de la production pharmaceutique (1958-2008)
- ✓ Une capacité de production très importante (43 millions unités de vente par an)
- ✓ Un savoir-faire élevé dans le domaine de la production, contrôle et analyse
- ✓ Une surface de stockage de 6.600 m² (4.600 palettes)

I.2.3 Organigramme du L'unité de Dar El Beida

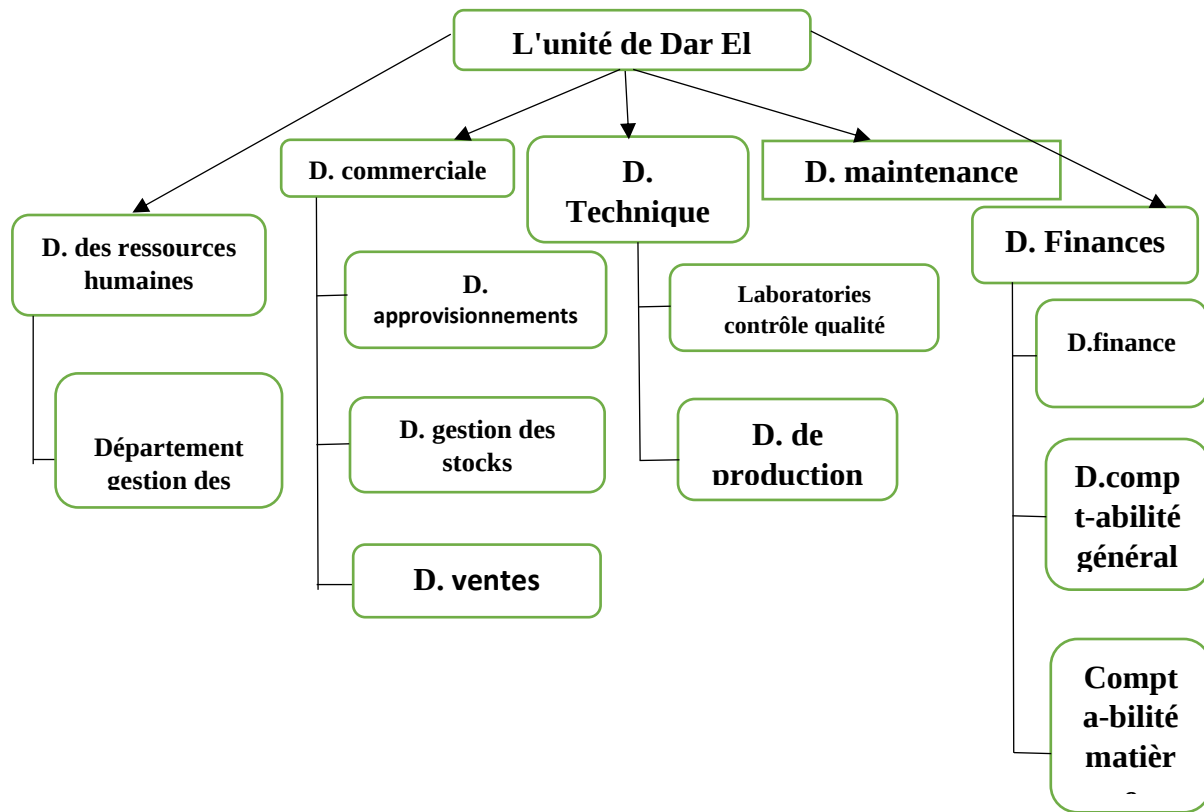


Figure 2 : les départements de L'unité de Dar El Beida

I.3 Motivation

L'évolution rapide de la technique de l'apprentissage profond, qui fournit une puissance de capacité d'analyse matérielle et de modélisation prédictive est l'objet de recherche. La science pharmaceutique, connue pour sa dynamique et sa sensibilité aux conditions du marché, nécessite des prédictions fiables fondées sur des données pour soutenir les besoins de matériel spécifiques pour augmenter l'efficacité de la production, réduire les gaspillages et garantir des délais de livraison. En prenant Saidal en tant qu'exemple spécifique, l'objectif est de résoudre un problème

de fonctionnement critique qui peut être abordé au moyen des meilleures pratiques en matière d'apprentissage profond. De plus, les valeurs et les connaissances tirées de cette recherche peuvent être appliquées à l'ensemble du secteur pharmaceutique, ce qui promet des avantages et des progrès considérables.

I.4 Problématique

Dans le cadre de notre projet de fin d'études, nous avons examiné le groupe industriel Saïdal (SPA), en particulier l'activité d'approvisionnement de matières premières de l'unité de Dar El Beida. La prévision des achats de matières premières constitue un enjeu crucial pour cette entreprise, car elle est essentielle pour maintenir l'équilibre entre la production et la consommation et ainsi éviter les pertes d'énergie et les surcoûts.

Le problème central que nous visons à résoudre est le manque important de méthodes systématiques et précises pour prédire les demandes du marché et les fluctuations saisonnières des besoins en matériaux à Saidal. Actuellement, ces prévisions sont principalement basées sur l'expérience et l'intuition du personnel de l'entreprise, ce qui peut entraîner une variabilité et un manque d'efficacité. Cette dépendance à l'égard des méthodes de prévision manuelles met en évidence une lacune critique que notre recherche vise à combler.

Notre objectif principal est de développer un cadre prédictif utilisant des algorithmes d'apprentissage profond pour améliorer la précision et la fiabilité des prévisions de la demande. En mettant en œuvre cette méthodologie avancée, nous visons à transformer l'approche de Saidal en matière de gestion des stocks et de planification de la production. Cette transformation devrait permettre d'améliorer l'efficacité opérationnelle, de mieux s'aligner sur la dynamique du marché et d'avoir une chaîne d'approvisionnement plus réactive.

Chapitre II : État de l'Art

Introduction

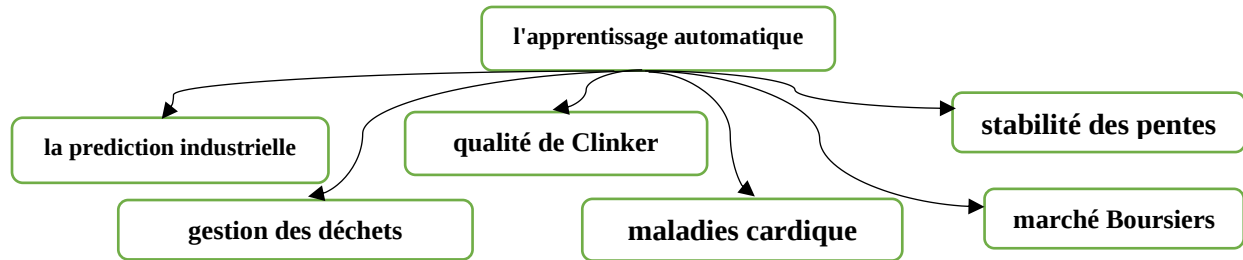


Figure 3 : les domaine d'application de l'apprentissage automatique

Les modèles tels que les réseaux neuronaux, les SVM, ainsi que les techniques d'optimisation comme PSO et IPSO sont appliqués avec succès pour améliorer la précision des prévisions, ouvrant de nouvelles perspectives dans l'application de l'intelligence artificielle à des défis concrets et diversifiés.

II.1 Synthèse bibliographique

Dzakiyullah, N. R., Hussin, B., Saleh, C., & Handani, A. M. [2] : Cette recherche compare les performances de deux modèles de prévision, les réseaux neuronaux (NN) et les machines à vecteurs de support (SVM), pour prédire les quantités de production dans le secteur de la fonderie. Utilisant des données privées de 2006 à 2013, les chercheurs ont divisé ces données en ensembles d'entraînement et de test via une validation croisée. Les paramètres des deux modèles ont été optimisés par essais et erreurs pour minimiser l'erreur de prévision. Les résultats montrent que le modèle NN a surpassé le modèle SVM, avec un RMSE (Root Mean Squared Error) de 1,713 en entraînement contre 1,718 pour SVM, et de 0,206 en test contre 0,309 pour SVM. Cette étude démontre la supériorité des réseaux neuronaux pour la prédiction des quantités de production dans l'industrie de la fonderie, soulignant leur avantage dans les applications de prévision industrielle.

Ayeleru, O. O., Fajimi, L. I., Oboirien, B. O., & Olubambi, P. A. [3]. : Cette étude se concentre sur la prévision des quantités de déchets solides municipaux (DSM) à Johannesburg, en Afrique du Sud, en utilisant des techniques d'apprentissage automatique. Confrontés à des données souvent indisponibles ou peu fiables, les chercheurs ont utilisé deux algorithmes : les réseaux neuronaux artificiels (ANN) et les machines à vecteurs de support (SVM). Les données historiques de Statistics South Africa ont servi de base pour des prévisions jusqu'en 2050. Les modèles ont été construits et testés dans MATLAB. Les résultats montrent que les modèles ANN, notamment avec une structure à 10 neurones (ANN10), ont atteint un coefficient de détermination (R^2) de 99,9 %, tandis que les modèles SVM linéaires ont obtenu un R^2 de 98,6 %. Ces techniques d'apprentissage automatique se révèlent efficaces pour prédire les DSM, aidant ainsi les décideurs à planifier la gestion des déchets de manière durable. La prévision avec ANN10 indique que Johannesburg

pourrait générer annuellement 1,95 million de tonnes de déchets solides d'ici 2050, avec une moyenne annuelle de 1,78 million de tonnes.

Boukhari, Y. (2020). [4] : Cette étude vise à prédire la perte de poids des matières premières lors de la production de clinker, un composant clé dans la fabrication du ciment. La perte de poids est cruciale pour la qualité du clinker mais difficile à évaluer précisément à cause de divers facteurs comme la composition, la taille des particules, la température et le temps de combustion. Les mesures expérimentales sont coûteuses et longues. Les chercheurs ont utilisé des modèles intelligents, dont le réseau neuronal artificiel optimisé par algorithme génétique (GA-ANN), les ensembles d'arbres de régression (RTE), les machines à vecteurs de support des moindres carrés (LS-SVM) et le système d'inférence neuro-floue adaptatif (ANFIS). Les résultats montrent que ces modèles prédisent efficacement la perte de poids, avec des coefficients de détermination ajustés (R^2) de 97,17 % à 99,31 %, et une erreur inférieure à 3,1 %. Cela démontre que les modèles d'apprentissage automatique peuvent remplacer les mesures expérimentales pour une prédiction rapide et économique, aidant ainsi les industriels à améliorer la qualité du ciment.

Singh, A., & Kumar, R. (2020), February [5] : Cette étude se focalise sur l'utilisation de l'apprentissage automatique pour prédire les maladies cardiaques, un défi crucial dans le domaine médical en raison de l'importance vitale du cœur. La précision maximale est essentielle pour le diagnostic et la prédiction des maladies cardiaques, car même de petites erreurs peuvent avoir des conséquences graves, voire fatales. Face à une augmentation exponentielle des décès liés aux problèmes cardiaques, il est impératif de développer des systèmes de prédiction fiables pour sensibiliser efficacement à ces maladies. L'apprentissage automatique, branche de l'intelligence artificielle, propose des solutions prometteuses pour anticiper ces événements à partir de données naturelles. Les chercheurs ont évalué la précision de plusieurs algorithmes d'apprentissage automatique comme k-nearest neighbor, arbre de décision, régression linéaire et machine à vecteur de support (SVM) dans la prédiction des maladies cardiaques, utilisant le jeu de données UCI repository pour l'entraînement et le test. Cette étude a été menée à l'aide du bloc-notes Anaconda (Jupyter), plateforme facilitant l'analyse précise grâce à ses bibliothèques Python et ses en-têtes. Les résultats obtenus fournissent des informations cruciales pour le développement de systèmes de prédiction des maladies cardiaques, contribuant ainsi à l'amélioration des soins de santé et potentiellement à la sauvegarde de vies supplémentaires.

Hegazy, O., Soliman, O. S., & Salam, M. A. [6] : Cette étude présente un modèle d'apprentissage automatique visant à prédire les prix du marché boursier, un sujet d'intérêt majeur pour maximiser les gains des investisseurs. Le modèle proposé combine l'optimisation par essaims de particules (PSO) avec la machine à vecteur de support des moindres carrés (LS-SVM). PSO est utilisé pour ajuster les paramètres du LS-SVM afin d'améliorer la précision de la prédiction en évitant le

surajustement et les minima locaux. Le modèle a été testé sur treize ensembles de données financières et comparé à un réseau de neurones artificiels utilisant l'algorithme de Levenberg-Marquardt (LM). Les résultats montrent que le modèle proposé offre une meilleure précision de prédiction, soulignant le potentiel de PSO dans l'optimisation du LS-SVM pour prédire les prix boursiers. Cette recherche apporte une contribution significative dans le domaine de la prédiction boursière en utilisant des techniques avancées d'apprentissage automatique, ouvrant ainsi de nouvelles perspectives pour les investisseurs cherchant à maximiser leurs gains sur les marchés financiers.

Wang, Y., Du, E., Yang, S., & Yu, L. [7] : Cette étude propose un modèle d'algorithme amélioré, combinant l'optimisation par essaims de particules (IPSO) et les machines à vecteurs de support (SVM), appelé IPSO-SVM, pour prédire la stabilité des pentes. Les méthodes traditionnelles d'évaluation de la stabilité des pentes sont limitées par leur dépendance à l'expérience, leurs temps de calcul longs et leur incapacité à bien gérer les caractéristiques floues et non linéaires des paramètres. Pour surmonter ces obstacles, le modèle IPSO-SVM a été testé sur 28 jeux de données d'ingénierie des pentes, incluant des paramètres tels que la densité apparente, la cohésion, l'angle de frottement, l'angle et la hauteur de pente, ainsi que le rapport de pression d'eau interstitielle. La stabilité des pentes, mesurée par le facteur de sécurité, est la variable de sortie. Les résultats montrent que le modèle IPSO-SVM surpasse nettement les modèles SVM et PSO-SVM traditionnels, avec une erreur quadratique moyenne jusqu'à 8,7 fois plus faible et une erreur relative moyenne de seulement 1,1%, la plus grande erreur étant de 1,3%. La dispersion des erreurs est très faible, autour de 1%. Ces résultats confirment l'efficacité et la fiabilité du modèle IPSO-SVM pour prédire avec précision la stabilité des pentes, ouvrant de nouvelles perspectives pour l'ingénierie pratique dans ce domaine.

Chaabene, W. B., Flah, M., & Nehdi, M. L. [8] : Cet article présente une étude sur l'utilisation des modèles d'apprentissage automatique pour prédire avec précision les propriétés mécaniques du béton. La prédiction de ces propriétés est une préoccupation importante, car elles sont souvent exigées par les codes de conception. L'émergence de nouveaux mélanges de béton et de nouvelles applications a motivé les chercheurs à trouver des modèles fiables pour prédire la résistance mécanique. Les modèles empiriques et statistiques, comme la régression linéaire et non linéaire, ont été largement utilisés, mais leur développement nécessite un travail expérimental laborieux et peuvent fournir des résultats inexacts lorsque les relations entre les propriétés du béton, la composition du mélange et les conditions de durcissement sont complexes.

Chapitre III : création d'un modèle de machine learning

Introduction

Ce chapitre se concentre sur la création d'un modèle de machine learning utilisant le Support Vector Machine (SVM) pour prédire la quantité de matière première nécessaire. Le processus est organisé en 12 étapes essentielles.

Ces étapes illustrent comment construire un pipeline complet pour la prédiction de quantités avec un modèle SVM, en assurant une gestion rigoureuse des données et une évaluation détaillée des performances du modèle.

III.1 Définition de l'intelligence artificielle

L'intelligence artificielle (IA) est une branche des sciences et technologies qui vise à reproduire le fonctionnement du cerveau humain par la machine. L'objectif est d'améliorer et d'étendre l'intelligence de la machine en permettant des raisonnements humains. Les progrès dans ce domaine ont été rapides, allant de la compréhension de textes à l'analyse d'images, et ont mené à la création de systèmes d'apprentissage qui ont donné naissance à de nouvelles disciplines. [9]

III.2 L'intelligence Artificielle et ses sous domaines

L'IA englobe tous les systèmes capables de simuler l'intelligence humaine. À l'intérieur de ce cercle, le machine Learning (apprentissage automatique) se concentre sur les algorithmes permettant aux machines d'apprendre à partir de données. Enfin, au cœur de cette hiérarchie, le deep Learning (apprentissage profond) représente une sous-catégorie du machine Learning, utilisant des réseaux de neurones profonds pour résoudre des problèmes complexes.

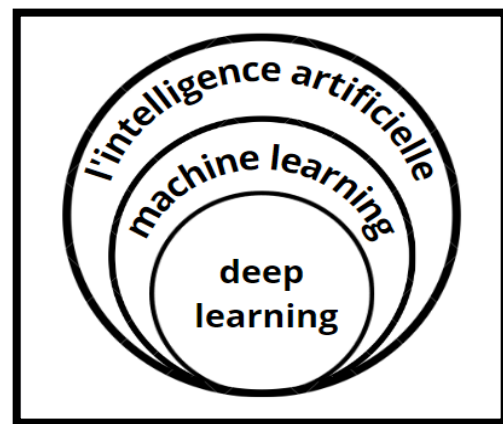


Figure 4: l'intelligence artificielle et ses sous domaines

III.3 Apprentissage automatique (Machine Learning)

Le "Machine Learning" est une méthode spécifique au sein du domaine plus vaste de l'intelligence artificielle. Son principe repose sur l'extraction d'informations à partir de l'analyse de données. Pendant ce processus d'apprentissage, la machine identifie des relations entre les différentes entrées et sorties. Ce procédé permet au système d'apprendre à partir des données et de s'améliorer au fil du temps en découvrant ces corrélations. [10]

III.4 Apprentissage profond (Deep Learning)

Cette discipline nommée Apprentissage profond est une sous discipline de l'Apprentissage automatique.

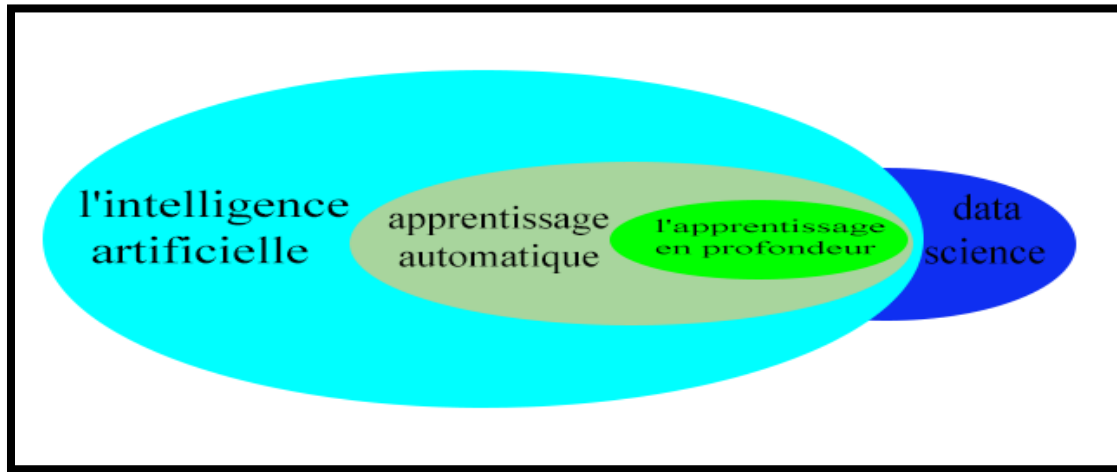


Figure 5: la relation entre Apprentissage automatique et Apprentissage profond

III.5 Les Types d'apprentissage automatique

L'apprentissage supervisé : entraîne un algorithme à associer des entrées à des sorties à partir de données étiquetées.

L'apprentissage semi-supervisé : combine des données étiquetées et non étiquetées pour améliorer la classification.

L'apprentissage non-supervisé : utilise des techniques comme le clustering et les K-means pour analyser des données sans entraînement préalable.

L'apprentissage par renforcement : permet aux agents d'interagir avec leur environnement pour maximiser les récompenses via une série d'actions.

III.6 Les techniques du d'apprentissage automatique utilisé dans ce chapitre

III.6.1 Support Vector Machine (SVM)

C'est l'un des algorithmes les plus populaires et les plus utilisés pour les tâches de classification. Cet algorithme d'apprentissage supervisé vise à trouver, dans un espace de n dimensions, la ligne qui sépare au mieux les données en différentes classes. Cette frontière linéaire, appelée hyperplan, permet de classer facilement les données dans la catégorie appropriée. [11]

III.6.2 SVR (Support Vector Regression)

SVR est une extension des SVM pour les tâches de régression. Contrairement à SVM qui est utilisé pour les tâches de classification, SVR essaie de trouver une fonction qui correspond aux données d'entraînement dans une tolérance d'erreur spécifiée.

SVR est utilisé dans des situations où il est nécessaire de prédire une valeur continue, comme la prévision des prix, les valeurs de température, etc.

Différence entre SVM et SVR : SVM Utilisé pour les tâches de classification et SVR Utilisé pour les tâches de régression.

III.6.2.1 Objectif de SVR

L'objectif de SVR est de trouver une fonction $f(x)$ qui prédit les valeurs cibles y avec des déviations dans une marge d'erreur ϵ . SVR minimise l'erreur et pénalise les prédictions en dehors de cette tolérance.

III.6.2.2 Fonctionnement de SVR

Le SVR vise à trouver une fonction prédictive qui minimise les erreurs tout en respectant une tolérance d'erreur ϵ . Il identifie les vecteurs de support, des points de données cruciaux influençant la prédiction. La méthode inclut une fonction de coût qui ignore les erreurs dans la zone de tolérance ϵ et pénalise celles en dehors, assurant une prédiction précise des valeurs continues.

III.7 Matériel utilisé

Table 1: Matériel utilisé pour création de model svm

	POSTE DE TRAVAIL
Pc	Lenovo ideapad
Système d'exploitation	Windows 10
Processeur	AMD Ryzen7 3700u with radeon vega mobile
RAM et Type de système	8 GB

III.8 Les Etapes de création d'un modèle de machine learning utilisant le (SVM) et (SVR)

III.8.1 Importation des Bibliothèques

```
import pandas as pd
import numpy as np
import joblib
from sklearn.preprocessing import LabelEncoder, StandardScaler
from sklearn.model_selection import train_test_split
from sklearn.svm import SVR
from sklearn.metrics import mean_squared_error
from sklearn.model_selection import cross_val_score
```

Figure 6:code de l'importation des bibliothèques nécessaire

Pour commencer, nous importons les bibliothèques essentielles suivant :

ETAPE 1 : manipule et analyse des données (Pandas).

ETAPE 2 : Facilite de la sérialisation et dé sérialisation d'objets, (Importons joblib)

ETAPE 3 : Convertir les valeurs catégorielles en valeurs numériques et normaliser les caractéristiques, (LabelEncoder et StandardScaler de Scikit-learn)

ETAPE 4 : Importée pour diviser les données en ensembles d'entraînement et de test. (La fonction train_test_split).

III.8.2 Charger un fichier CSV dans un DataFrame Pandas

```
# Load the dataset
df = pd.read_csv('dataset/NewDaya2017.csv')
```

Figure 7: code de chargement de fichier dataset newdaya2017.csv

ETAPE 1 : Charge le fichier CSV dans un DataFrame Pandas, nous utilisons la bibliothèque (Pandas avec l'alias pd).

ETAPE 2 : Lire le fichier CSV et de le charger dans un DataFrame.

ETAPE 3 : Stocker le résultat dans une variable nommée df.

III.8.3 Encoder des variables catégorielles en valeurs numériques l'aide de LabelEncoder

```
# Encode categorical variables
label_encoder_desigart = LabelEncoder()
label_encoder_nom = LabelEncoder()

df['DESIGART'] = label_encoder_desigart.fit_transform(df['DESIGART'])
df['NOM'] = label_encoder_nom.fit_transform(df['NOM'])
```

Figure 8: code de l'encodage des variables

Ce processus est couramment utilisé dans le prétraitement des données pour les modèles de machine Learning qui ne peuvent pas traiter des variables catégorielles directement.

ETAPE 1 : Encoder les variables catégorielles on crée deux instances de LabelEncoder, (la colonne DESIGART et la colonne NOM).

ETAPE 2 : Apprendre les classes uniques dans chaque colonne et convertir ces classes en valeurs numériques à l'aide de la fonction fit_transform.

III.8.4 Définir les features (caractéristiques) et la target (cible) dans un DataFrame Pandas

Cette préparation des données est une étape essentielle avant de pouvoir les utiliser pour entraîner un modèle prédictif.

```
# Define features and target
X = df[['DESIGART', 'Valeur', 'NOM']]
y = df['QTE_ENT']
```

Figure 9 : code de préparation des données

ETAPE 1 : sélectionne les colonnes dans le DataFrame df

ETAPE 2 : stocker respectivement dans les variables X et y, (X représente les données d'entrée et y les valeurs de sortie).

Exemple, les colonnes 'DESIGART', 'Valeur', et 'NOM' sont choisies comme caractéristiques. Le target est la variable dépendante ou la variable à prédire, ici la colonne 'QTE_ENT'.

III.8.5 Normaliser les caractéristiques d'un DataFrame à l'aide de StandardScaler

```
# Scale features
scaler = StandardScaler()
X = scaler.fit_transform(X)
```

Figure 10 : code de normalisation des données

Normaliser les Données” signifie ajuster les valeurs mesurées sur différentes échelles à une notion commune,

ETAPE 1 : calcule la moyenne et l'écart-type de chaque caractéristique dans X

ETAPE 2 : transforme les données en soustrayant la moyenne et en divisant par l'écart-type à l'aide de fonction fit_transform(X))

ETAPE 3 : remplacées les valeurs de X par des valeurs normalisées.

III.8.6 Diviser des tableaux ou les matrices

```
# Split the data into training and testing sets
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2,
```

Figure 11 : diviser les données

ETAPE 1 : Diviser les matrices en sous-ensembles.

ETAPE 2 : Tester de manière aléatoire à l'aide de la fonction fit_transform(X) avec Test_size=0.2.

```
# Train the SVM model
svm_model = SVR()
svm_model.fit(X_train, y_train)
```

▼ SVR

SVR()

Figure 12 : code d'entraînement des données

Pour entraîner le modèle, nous utilisons la méthode **fit** avec les données d'entraînement **X_train** (caractéristiques) et **y_train** (cible). Cette méthode ajuste les paramètres du modèle pour qu'il s'adapte au mieux aux données d'entraînement.

III.8.7 Évaluer un modèle de régression SVM

```
# Evaluate the model
y_pred = svm_model.predict(X_test)
print(f'Mean Squared Error: {mean_squared_error(y_test, y_pred)}')
```

Figure 13 : code d'évaluation des données

ETAPE 1 : Prédire les valeurs de la variable cible (y) en utilisant les données de test (**X_test**).

ETAPE 2 : Stocker dans **y_pred** contient les valeurs prédites.

ETAPE 3 : Calcule l'erreur quadratique moyenne (MSE) entre les valeurs réelles de la variable cible (**y_test**) et les valeurs prédites (**y_pred**).

III.8.8 Validation croisée sur un modèle SVM (Support Vector Machine)

```
# Perform cross-validation
cv_scores = cross_val_score(svm_model, X, y, cv=5, scoring='neg_mean_squa
```

Figure 14 : code d'évaluation croisée

ETAPE 1 : effectuer une validation croisée 5 fois à l'aide de la fonction **cross_val_score**

ETAPE 2 : stocker les résultats dans le tableau **cv_scores**.

Ce code Python utilise la bibliothèque **scikit-learn** pour effectuer une validation croisée sur un modèle SVM

L'erreur quadratique moyenne négative est utilisée comme métrique d'évaluation.

III.8.9 Convertit les scores de validation croisée en pourcentages

```
# Convert mean squared error to percentage
cv_scores_percentage = np.sqrt(-cv_scores) / np.mean(y) * 100
```

Figure 15 : code de convertir en pourcentages

Ce code Python prend les scores de validation croisée obtenus précédemment (qui sont des valeurs négatives de l'erreur quadratique moyenne) et les

ETAPE 1 : Sélectionner les valeurs négatives de l'erreur quadratique moyenne

ETAPE 2 : Calculer la racine carrée des valeurs négatives de l'erreur quadratique moyenne,

ETAPE 3 : Divise les résultats par la moyenne des valeurs cibles.

ETAPE 4 : Multiplie le résultat par 100 pour obtenir un pourcentage.

Ce code permet de comprendre la performance du modèle en termes de pourcentage de la moyenne des cibles, ce qui est souvent plus interprétable que l'erreur quadratique moyenne brute.

III.8.10 Calcule la moyenne et l'écart type des erreurs en pourcentage obtenues précédemment

```
# Calculate mean and standard deviation of the scores
mean_score = np.mean(cv_scores_percentage)
std_dev_score = np.std(cv_scores_percentage)

print(f'Mean Percentage Error: {mean_score:.2f}%')
print(f'Standard Deviation of Percentage Error: {std_dev_score:.2f}%')
```

Figure 16 : code de Calcul de la moyenne et l'écart type

Ce code calcule et affiche la moyenne et l'écart type des erreurs en pourcentage des scores de validation croisée, ce qui permet de comprendre à quel point les erreurs sont centrées et dispersées autour de la moyenne.

ETAPE 1 : Calculer la moyenne des valeurs dans le tableau à l'aide de la fonction `Np.Mean(cv_scores_percentage)`

ETAPE 2 : calculer l'écart type des valeurs dans le tableau `cv_scores_percentage` à l'aide de la fonction `np.std(cv_scores_percentage)`

ETAPE 3 : Formater et afficher la moyenne des erreurs en pourcentage avec deux chiffres après la virgule à l'aide de la fonction `Print (f'Mean Parentage Error : {mean_score :2f} %')`

ETAPE 4 : Afficher les résultats en pourcentage.

Print (f'Standard Deviation of Percentage Error : {std_dev_score :2f} %') : Cette ligne utilise une f-string pour formater et afficher l'écart type des erreurs en pourcentage avec deux chiffres après la virgule. Le résultat est également affiché en pourcentage.

III.8.11 Sauvegarder plusieurs objets liés à un modèle d'apprentissage automatique.

```
# Save the trained SVM model
joblib.dump(svm_model, 'svm_model.joblib')

# Save the label encoders and scaler
joblib.dump(label_encoder_desigart, 'label_encoder_desigart.joblib')
joblib.dump(label_encoder_nom, 'label_encoder_nom.joblib')
joblib.dump(scaler, 'scaler.joblib')

['scaler.joblib']
```

Figure 17 : code de sauvegarder de model

ETAPE 1 : sauvegarde un modèle SVM entraîné dans un fichier nommé svm_model. Joblib à l'aide de la fonction joblib.

ETAPE 2 : Sauvegarde un encodeur de labels a l'aide de fonction Joblib. Dump

ETAPE 3 : Sauvegarde un objet de mise à l'échelle des données appelé scaler dans un fichier nommé scaler. Joblib.

Conclusion

Dans ce chapitre nous avons présenté la démarche à suivre en détaille pour l'implantation d'un modèle de prédiction d'approvisionnement en quantité de médicament au niveau de groupement Sidal – Dar Baida –Algerie.

Les résultats obtenus à partir des simulations sur le logiciel python en utilisant la méthode SVM seront discutée dans le chapitre suivant

Chapitre IV : interprétation et discussion des résultats

Introduction

Dans ce chapitre, nous abordons la présentation des résultats obtenus par notre modèle de prédiction d'approvisionnement en quantité des matières premières au niveau de groupement Saidal DAR Beida Algérie. Nous commencerons par la présentation des résultats, mettant en évidence les caractéristiques de chaque produit. Ensuite, nous interprétons des résultats, en analysant les performances de chaque produit. Enfin, nous discutons les courbes associées à ces données, en traçant d'une part les variations de la quantité approvisionnée en fonction du temps et d'autre par la prédiction de l'horizon suivant.

IV.1 Présentation d'interface

Dans le cadre de notre projet de prédiction de quantité de matières premières en utilisant les machines à vecteurs de support (SVM), nous avons développé une interface utilisateur en Python. Cette interface met l'accent sur l'importance d'un système de prédiction performant pour anticiper les besoins et prévenir les quantités de matières premières.

- ✚ **Objectifs de l'interface :** Dans le contexte de votre projet, l'interface SVM permettra d'anticiper les besoins en matières premières et de prévenir les quantités nécessaires.
- ✚ **Développement de l'Interface**
- ✚ **Configuration l'environnement de développement :** il est essentiel de vérifier que les conditions préalables techniques sont remplies :
- ✚ **Développement des Modules :** Nous présentons les dossiers spécifiques pour les scripts Python (voir Fig. 1), les fichiers de base de données et les ressources (images, fichiers de configuration).

Nom	Modifié le	Type	Taille
dataset	06/19/2024 11:43	Dossier de fichiers	
Export	06/14/2024 02:00	Dossier de fichiers	
modelINN	06/14/2024 01:19	Dossier de fichiers	
modelSVM	06/09/2024 02:32	Dossier de fichiers	
interfaceSVM.py	06/09/2024 02:38	Fichier PY	3 Ko
Svmmodel.ipynb	06/09/2024 02:54	Fichier IPYNB	10 Ko

Figure 18 : Organisation du contenu

IV.2 Fonctionnement de l'interface

Pour mettre en œuvre l'interface de prédiction des quantités de matières premières, nous avons exécuté le script Python "interfaceSVM.py" à partir de l'invite de commandes Windows

PowerShell. Cette approche permet d'exécuter le programme et d'afficher les messages de débogage et d'avertissement liés à l'utilisation de la bibliothèque TensorFlow, comme illustré dans la figure ci-dessous.

```

PS C:\Users\DELL\OneDrive\Bureau\PFE2024\RNN nd SVM\DayaModels (1)\DayaModels> python interfaceSVM.py
C:\Users\DELL\AppData\Local\Programs\Python\Python312\Lib\site-packages\sklearn\base.py:376: InconsistentVersionWarning: Trying to unpickle estimator SVR from version 1.2.2 when using version 1.5.0. This might lead to breaking code or invalid results. Use at your own risk. For more info please refer to: https://scikit-learn.org/stable/model_persistence.html#security-maintainability-limitations
warnings.warn(
C:\Users\DELL\AppData\Local\Programs\Python\Python312\Lib\site-packages\sklearn\base.py:376: InconsistentVersionWarning: Trying to unpickle estimator LabelEncoder from version 1.2.2 when using version 1.5.0. This might lead to breaking code or invalid results. Use at your own risk. For more info please refer to: https://scikit-learn.org/stable/model_persistence.html#security-maintainability-limitations
warnings.warn(
C:\Users\DELL\AppData\Local\Programs\Python\Python312\Lib\site-packages\sklearn\base.py:376: InconsistentVersionWarning: Trying to unpickle estimator StandardScaler from version 1.2.2 when using version 1.5.0. This might lead to breaking code or invalid results. Use at your own risk. For more info please refer to: https://scikit-learn.org/stable/model_persistence.html#security-maintainability-limitations
warnings.warn(

```

Figure 19 : exécuté le script Python

Figure 20 : interface model

IV.3 Discussion Et Comparaison Des Résultat de prédiction

Table 2 : Résultat de prédiction sur 2023

N°	PRODUIT	FOURNISSEUR	QUANTITE kg (2023)	VALEUR DZD	QUANTITE PRÉDICTIF kg (2023)
1	PROSPECTUS PARALGAN 500MG 10.5CMX15CM	EMBAG BBA	277,00	66528.00	171,71
2	ETUIS RHUMAFED COMPRIMES	EDIALCART	460,00	588806.40	172,55

3	LAURYL SULFATE DE SODIUM POUDRE	AXOPHARMA	175,00	427498.008	486.30
4	SULPIRIDE POUDRE	MEFAM LAB LIMITED	770,00	4176908.977	1800,65
5	PANTHOTENATE DE CALCIUM	JAB COMPANY	105,00	1110375.000	193,38

Le tableau présente une comparaison entre les valeurs en temps réel et les valeurs prédites par notre modèle.

Comme le montre le tableau ci-dessus, le modèle surestime largement la plupart des matières premières, ce qui va à l'encontre de l'objectif même du modèle. Cette estimation peut entraîner un sur stockage et une augmentation des coûts de détention, alors que certaines quantités prédites sont proches de la valeur réelle des matières premières,

Nous pouvons constater que le modèle proposé n'est pas fiable pour effectuer une prédiction correcte et que la marge d'erreur ne peut être tolérée, cela peut être lié à plusieurs raisons dans l'architecture du modèle, pour éviter cela, nous devons essayer un noyau différent.

Le choix du noyau dans les machines à support vectoriel (SVM) dépend largement de l'ensemble de données ; différents noyaux peuvent donner de bons résultats dans notre cas :

- ✚ **Noyau linéaire** : Elle est utile pour les données linéairement séparables, ce qui signifie qu'elles peuvent être séparées par une simple ligne (en 2D) ou un hyperplan (en dimensions supérieures). Il est plus simple et plus efficace en termes de calcul, mais il peut ne pas donner de bons résultats avec des ensembles de données complexes.
- ✚ **Noyau polynomial** : Un noyau polynomial est une forme plus généralisée du noyau linéaire. Le noyau polynomial peut distinguer un espace d'entrée courbe ou non linéaire. Cela permet d'obtenir des limites plus complexes, ce qui peut être utile pour les ensembles de données qui présentent une relation non linéaire mais néanmoins assez systématique.
- ✚ **Noyau sigmoïde** : Le noyau sigmoïde possède des propriétés très similaires aux noyaux RBF et polynomial et est souvent utilisé dans la pratique en raison de son origine dans les réseaux neuronaux. Le noyau sigmoïde transforme les données dans un espace où elles deviennent linéairement séparables.
- ✚ **Noyau RBF** : Le noyau de la fonction de base radiale est une fonction de noyau populaire couramment utilisée dans la classification SVM.

Interprétation

- nous constatons qu'il est plus approprié d'exclure les noyaux polynomiaux linéaires et sigmoïdes en raison du biais élevé que présente le noyau linéaire et de son hypothèse selon laquelle les données sont linéairement séparables, de la complexité de calcul du noyau polynomial et des problèmes de sur ajustement qui entraînent une mauvaise généralisation à des données non vues,
- le noyau sigmoïde n'est pas couramment utilisé car, contrairement aux trois autres noyaux, les noyaux sigmoïdes ne sont pas toujours définis positifs, ce qui peut poser des problèmes lors de l'apprentissage. C'est pourquoi les résultats étaient extrêmement éloignés des valeurs réelles dans le premier tableau : nous avons rencontré le problème de l'ajustement excessif et le modèle n'a pas bien géré les données inédites.

Le meilleur noyau est le noyau RBF en raison de sa flexibilité et de sa moindre sensibilité au choix des paramètres, ce qui convient parfaitement à l'ensemble de données complexe que nous avons entre les mains et au manque de paramètres avec lesquels travailler.

Après avoir travaillé avec le noyau RBF, les résultats ont été conformes aux attentes et nous avons constaté une amélioration considérable des performances du modèle.

Table 3 : Résultat de prédiction sur 2023 Après Avoir Travaillé Avec Le Noyau RBF

N°	PRODUIT	FOURNISSEUR	QUANTITE KG (2023)	VALEUR DZD	QUANTITE PRÉDICTIF KG (2023)
1	PROSPECTUS PARALGAN 500MG 10.5CMX15CM	EMBAG BBA	277,00	66528,00	261,63
2	ETUIS RHUMAFED COMPRIMES	EDIALCART	460,00	588806.40	396,73
3	LAURYL SULFATE DE SODIUM POUDRE	AXOPHARMA	175,00	427498,008	174.70
4	SULPIRIDE POUDRE	MEFAM LAB LIMITED	770,00	4176908,977	735,27
5	PANTHOTENATE DE CALCIUM	JAB COMPANY	105,00	1110375,000	104,18

Dans le tableau, nous avons comparé les quantités réelles d'un exemple de 5 produits en 2023 avec les quantités prédites par le modèle SVM après avoir changé le noyau en noyau RBF.

- ❖ Pour le produit "Prospectus Paralgan 500mg 10.5cmx15cm", la quantité réelle utilisée en 2023 était de 277 unités, alors que le modèle SVM a prédit une quantité de 261,63 unités. Il en résulte une différence mineure de 15,37 unités, ce qui indique que la prédiction du modèle est assez précise.
- ❖ Dans le cas de "Etuïs Rhumafed Comprimés", la quantité réelle utilisée était de 460 unités, mais le modèle a prédit une quantité légèrement inférieure de 396,73 unités. La différence est donc plus importante (63,27 unités), mais elle reste dans une fourchette raisonnable et peut être tolérée
- ❖ Pour le "Lauryl Sulfate de Sodium Poudre", la quantité réelle utilisée était de 175 unités, et le modèle a prédit une quantité très proche de 174,70 unités. Il en résulte une différence extrêmement faible de seulement 0,30 unité, ce qui démontre l'excellente performance prédictive du modèle.
- ❖ Le modèle a également obtenu de bons résultats pour le "Sulpiride Poudre", avec une quantité réelle de 770 unités et une quantité prédite de 735,27 unités. La différence est de 34,73 unités.
- ❖ Enfin, pour le "Panthotenate de Calcium" fourni par JAB COMPANY, la quantité réelle était de 105 unités, et le modèle a prédit une quantité de 104,65 unités. Cela montre une prédiction très précise avec une différence négligeable de seulement 0,35 unité.

Dans l'ensemble, après avoir changé de noyau, le modèle présente d'excellents résultats et des performances améliorées, comme indiqué précédemment, par rapport au module précédent utilisant les autres noyaux.

IV.4 La présentation des résultats du model sur 2024

Table 4 : Résultat de prédiction sur 2024

N°	PRODUIT	FOURNISSEUR	QUANTITE KG (2023)	VALEUR DZD	QUANTITE KG (2024)
1	PROSPECTUS PARALGAN 500MG 10.5CMX15CM	EMBAG BBA	277,00	66528,00	171.71
2	ETUIS RHUMAFED COMPRIMES	EDIALCART	460,00	588806.40	370,65

3	LAURYL SULFATE DE SODIUM POUDRE	AXOPHARMA	175,00	427498,008	172.30
4	SULPIRIDE POUDRE	MEFAM LAB LIMITED	770,00	4176908,977	882,95
5	PANTHOTENATE DE CALCIUM	JAB COMPANY	105,00	1110375,000	132,10
6	MANNITOL	MEFAM LAB LIMITED	420,00	335522,376	369,12
7	TUBE CLOGEL 30GR	EL MAADEN	381,180	3503046,584	455.64
8	FILM THERMAL ROUGE 95MM	ETIQUAL SARL	140,00	238000	172
9	STEARATE DE MAGNESIUM	MEFAM LAB LIMITED	693,00	228360,7536	371.99
10	PANTHOTENATE DE CALCIUM	JAB COMPANY	105,00	1110375,00	173.38

Le tableau comparant les quantités de matières premières demandées en 2023 avec les quantités prévues pour 2024 montre des tendances variables pour différents produits.

- ❖ Pour Prospectus Paralgan 500mg 10.5cmx15cm, la demande devrait diminuer de manière significative, passant de 277 unités à 171,71 unités, soit une baisse de 105,29 unités.
- ❖ Etais Rhumafed Comprimés devrait également connaître une réduction, passant de 460 unités à 370,65 unités, soit une baisse de 89,35 unités.
- ❖ Le Lauryl Sulfate de Sodium Poudre montre une légère diminution de 175 unités à 172,30 unités, soit une réduction mineure de 2,70 unités.
- ❖ En revanche, le Sulpiride Poudre devrait passer de 770 unités à 822,95 unités, soit une augmentation de 52,95 unités.
- ❖ Le Panthotenate de Calcium devrait passer de 105 unités à 132,10 unités, soit une augmentation de 27,10 unités.
- ❖ Le mannitol devrait passer de 420 unités à 369,12 unités, soit une réduction de 50,88 unités.
- ❖ Le tube Clogel 30gr montre une augmentation significative de 381,18 unités à 455,64 unités, soit une augmentation de 74,46 unités.

- ❖ Le film Thermal Rouge 95mm devrait passer de 140 à 172 unités, soit une augmentation de 32 unités. Le Stéarate de Magnésium devrait connaître une baisse substantielle, passant de 693 unités à 371,99 unités, soit une chute de 321,01 unités.
- ❖ Enfin, la double entrée pour Panthotenate de Calcium montre une augmentation de 105 unités à 173,38 unités, soit une augmentation de 68,38 unités.

Dans l'ensemble, les quantités prévues pour 2024 indiquent à la fois des augmentations et des diminutions significatives de la demande pour différents produits, l'évolution de la demande étant tout à fait raisonnable.

IV.5 Représentation graphique des résultats

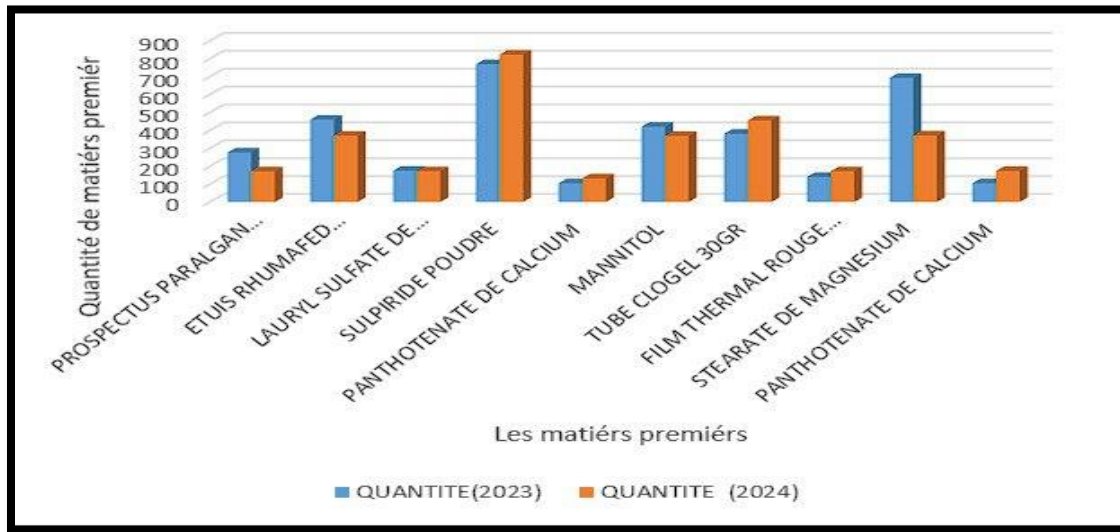


Figure 21 : Représentation de l'évolution des déferant matière entre 2023 et la réduction en 2024

Le graphe compare les quantités de diverses matières premières entre 2023 et 2024. Pour le PROSPECTUS PARALGAN, on observe une diminution notable en 2024 par rapport à 2023, passant de 350 à 250 unités. Les quantités d'ETUIS RHUMAFED augmentent légèrement, de 100 à 150 unités. Le LAURYL SULFATE DE... et le FILM THERMAL ROUGE connaissent une légère diminution, tandis que le SULPIRIDE POUDRE et le PANTHOTENATE DE CALCIUM (deuxième occurrence) voient une diminution plus marquée. À l'inverse, le MANNITOL, le TUBE CLOGEL 30GR et le STEARATE DE MAGNESIUM montrent une augmentation modérée, et le PANTHOTENATE DE CALCIUM (première occurrence) enregistre une augmentation notable. Les variations significatives dans les quantités de PANTHOTENATE DE CALCIUM et de SULPIRIDE POUDRE suggèrent une réévaluation des besoins ou une modification dans la production ou la demande de ces matières premières.

IV.6 Représentation graphique de l'évaluation de la performance du modèle SVM

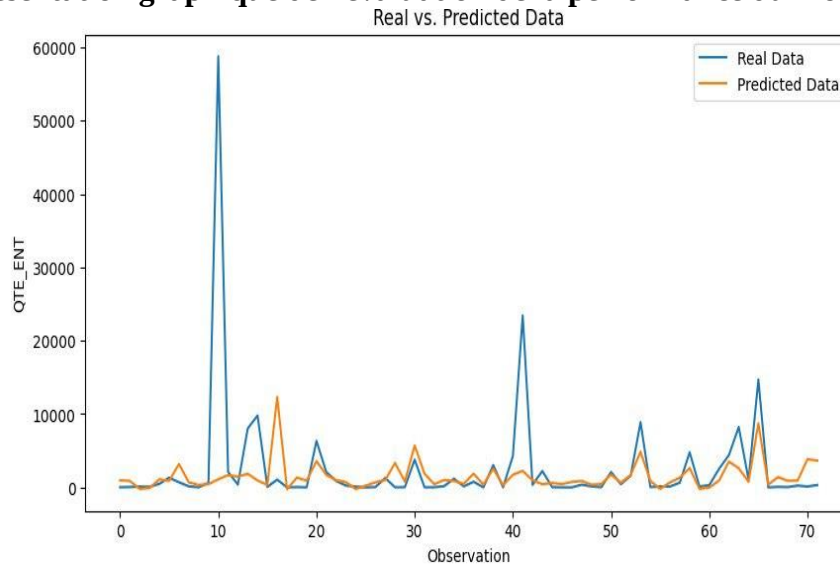


Figure 22: l'évaluation de la performance du modèle SVM

Le graphe compare les données réelles et les prédictions d'un modèle SVM, montrant que le modèle suit généralement les tendances des données réelles mais sous-estime les valeurs extrêmes et les pics marqués. Bien que les prédictions soient précises pour les valeurs faibles à modérées, les écarts deviennent significatifs pour les événements rares ou les anomalies. Cela suggère que le modèle a des difficultés à capturer les fluctuations importantes. Pour améliorer les performances, il pourrait être utile d'ajuster les paramètres du modèle ou d'ajouter plus de données d'entraînement, en particulier des exemples de valeurs extrêmes.

IV.7 Discussion des choix effectués

Lors de la mise en œuvre du modèle de prédiction, nous avons réalisé une étude comparative entre deux techniques d'intelligence artificielle : les machines à vecteurs de support (SVM) et les réseaux de neurones artificiels. Cette étude a été menée sur une même base de données, et le critère de comparaison était l'exactitude.

Définition de l'exactitude : désigne le pourcentage de prédictions correctes par rapport à l'ensemble des prédictions. Elle est calculée en utilisant quatre critères de probabilité

- ✚ Les vrais positifs (True Positif),
- ✚ Les vrais négatifs (True Negatif),
- ✚ Les faux positifs (False Positif)
- ✚ Les faux négatifs (False Negatif).

L'exactitude est influencée par la précision des valeurs positives (True et False) : plus elles sont élevées, plus l'exactitude de l'algorithme est élevée, tandis que des valeurs négatives (True et

False) plus élevées réduisent l'exactitude. L'expression mathématique pour calculer l'accuracy est la suivante :

$$\text{l'accuracy (exactitude)} = (TP + NP) / (TP + NP + TN + NN)$$

Les résultats de la comparaison sont les suivants :

Table 5 : Table de comparaison des techniques en fonction de l'accuracy

	Techniques	Accuracy
1	Artificial Neural Network	97,59
2	Support Vector Machine	78,59

Conclusion

En conclusion, l'évaluation de la performance du modèle SVM montre une capacité satisfaisante à prédire les tendances générales des données réelles, avec une bonne précision pour les valeurs faibles à modérées. Cependant, le modèle présente des limitations importantes lorsqu'il s'agit de prédire des valeurs extrêmes et des pics marqués. Ces écarts soulignent la nécessité d'optimiser d'avantage le modèle, soit en ajustant ses paramètres, soit en enrichissant l'ensemble de données d'entraînement avec des exemples de valeurs extrêmes. Cette analyse met en lumière l'importance de la robustesse du modèle face aux variations et aux anomalies des données pour améliorer sa fiabilité et sa précision.

Conclusion générale

Supply side est un défi majeur pour l'industrie, et pour saidal, compte tenu de la grande production et de la complexité des produits, c'est encore plus difficile.

Traditionnellement, cette gestion a été basée sur l'expérience individuelle, et des méthodes dépassées mais pour suivre cette grande transformation il est temps d'adopter une approche plus rigoureuse basée sur l'intelligence artificielle (IA). Notre projet PFE a porté sur l'exploitation d'une puissante méthode d'IA utilisée dans la prédiction de l'approvisionnement de matières premières à l'aide de modèles d'apprentissage automatique, en mettant l'accent sur les machines à vecteurs de support (SVM) et enfin sur une comparaison entre cette méthode et la méthode RNN utilisée sur le même ensemble de données.

En revisitant notre problématique, nous avons cherché à répondre à la question suivante : comment anticiper avec précision la demande future de matières premières ?

Le deuxième chapitre propose une exploration approfondie de l'état de l'art en matière de machines à vecteurs de support (SVM). Les SVM sont reconnus pour leur robustesse dans les tâches de classification et de régression, en particulier dans le traitement de données complexes et de haute dimension. Nous avons également examiné les applications des SVM dans divers secteurs, démontrant ainsi leur polyvalence et leur efficacité.

Nos recherches ont montré que les SVM, comparés à d'autres méthodes d'intelligence artificielle, constituent une alternative solide à l'approche traditionnelle. Leur capacité à traiter des données complexes et à généraliser à partir d'échantillons limités, ainsi que leur flexibilité, en font de puissants outils de prédiction.

Dans le troisième chapitre, nous avons détaillé les étapes de l'élaboration de notre modèle prédictif basé sur les SVM. Nous avons commencé par les définitions nécessaires et le contexte théorique, afin d'assurer une bonne compréhension des méthodologies employées. La configuration du modèle comprenait le prétraitement des données, la sélection des caractéristiques et le choix des fonctions de noyau appropriées. Nous avons expérimenté diverses configurations, en optimisant les hyper paramètres pour obtenir les meilleures performances. Nous avons également discuté des mesures d'évaluation utilisées pour évaluer la précision et la fiabilité du modèle, en veillant à ce que les prédictions soient robustes et exploitables.

Le quatrième chapitre présente les résultats de notre modèle SVM. En comparant les prédictions du modèle aux données en temps réel, nous avons démontré sa précision et sa fiabilité dans l'anticipation des besoins futurs en matières premières. Nous avons également comparé les performances de notre modèle SVM à celles d'un modèle RNN entraîné sur le même ensemble de données. La comparaison a mis en évidence les forces et les faiblesses de chaque approche, les SVM montrant des performances supérieures dans certains aspects en raison de leur capacité à bien généraliser à partir d'échantillons de données limités.

Nos recherches ont montré que les SVM offrent une alternative solide aux approches traditionnelles pour prédire les besoins en matières premières. Leur capacité à traiter des données complexes, associée à leur flexibilité, en fait de puissants outils de prédiction. Nous avons exploré diverses architectures et fonctions de noyau afin de maximiser le potentiel du modèle, pour finalement parvenir à un modèle prédictif efficace basé sur sept années de données historiques.

Cependant, notre étude ne se limite pas aux résultats actuels. Nous envisageons un avenir où la gestion des matières premières continuera d'évoluer avec les progrès de l'IA. Le développement continu des technologies de l'IA promet de nouvelles opportunités d'optimisation. De futurs modèles, potentiellement plus efficaces que les modèles actuels, pourraient voir le jour, révolutionnant nos pratiques et offrant des gains d'efficacité encore plus importants.

En conclusion, l'IA est une ressource précieuse pour l'optimisation de supply side. Elle permet de prendre des décisions plus éclairées, de réduire les coûts et d'anticiper les besoins futurs. À mesure que la technologie progresse, nous devons rester ouverts à l'innovation et à l'exploration de nouvelles méthodes. Approvisionnement est un domaine en constante évolution, et l'IA, en particulier grâce à des méthodes telles que les SVM, nous guidera vers un avenir plus efficace et plus durable.

En adoptant l'IA et en affinant continuellement nos approches, nous pouvons faire en sorte que des industries comme Sidal soient bien équipées pour relever les défis des demandes de production modernes, ce qui se traduira en fin de compte par une amélioration de l'efficacité opérationnelle et de la rentabilité.

Bibliographie

- [1] Belimane, W. (2010). Rapport de stage : cas Pharmal, filiale du groupe pharmaceutique Saidal en Algérie. Ecole des Hautes Etudes Commerciales. Disponible sur https://www.memoireonline.com/02/12/5367/m_Rapport-de-stage-cas-Pharmal-filiale-du-groupe-pharmaceutique-Saidal-en-Algerie.html
- [2] Dzakiyullah, N. R., Hussin, B., Saleh, C., & Handani, A. M. (2014). Comparison neural network and support vector machine for production quantity prediction. *Advanced Science Letters*, 20(10-11), 2129-2133.
- [3] Ayeleru, O. O., Fajimi, L. I., Oboirien, B. O., & Olubambi, P. A. (2021). Forecasting municipal solid waste quantity using artificial neural network and supported vector machine techniques: A case study of Johannesburg, South Africa. *Journal of Cleaner Production*, 289, 125671.
- [4] Boukhari, Y. (2020). Using Intelligent Models to Predict Weight Loss of Raw Materials During Cement Clinker Production. *Revue d'Intelligence Artificielle*,
- [5] Singh, A., & Kumar, R. (2020, February). Heart disease prediction using machine learning algorithms. In 2020 international conference on electrical and electronics engineering (ICE3) (pp. 452-457). IEEE.
- [6] Hegazy, O., Soliman, O. S., & Salam, M. A. (2014). A machine learning model for stock market prediction. *arXiv preprint arXiv:1402.7351*.
- [7] Wang, Y., Du, E., Yang, S., & Yu, L. (2022). Prediction and Analysis of Slope Stability Based on IPSO-SVM Machine Learning Model. *Geofluids*, 2022(1), 8529026.
- [8] Chaabene, W. B., Flah, M., & Nehdi, M. L. (2020). Machine learning prediction of mechanical properties of concrete: Critical review. *Construction and Building Materials*, 260, 119889.
- [9] : O. Pallanca et J. Read, « Principes Généraux et Définitions En Intelligence Artificielle », *Archives des Maladies du Cœur et des Vaisseaux - Pratique*, t. 2021, 1er déc. 2020. doi : 10.1016/j.amcp.2020.11.002.
- [10] K. Shailaja, B. Seetharamulu et M. A. Jabbar, « Machine Learning in Healthcare: A Review », in 2018 Second International Conference on Electronics, Communication and Aerospace Technology (ICECA), mars 2018, p. 910-914. doi : 10.1109/ICECA.2018.8474918.
- [11] Support Vector Machine (SVM) Algorithm - Javatpoint ». (), adresse : <https://www.javatpoint.com/machine-learning-support-vector-machine-algorithm> (visité le 04/07/2023).