

REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE

Ministère de l'enseignement supérieur et de la recherche scientifique



Université Saad Dahleb blida 1

Faculté des Sciences

Département d'Informatique

Projet de fin d'études en vue de l'obtention du

Diplôme de Master II en Systèmes informatiques et réseaux

**Mesure de la gravité des signes de dépression à partir des
réseaux sociaux**

Présenté par

Kritli Mohamed Chaouki

Boukenaoui Anfel

Soutenu le 10 septembre 2020

Devant le jury :

Présidente :	Hireche Celia	MAB	Université Saad dahleb
Examinatrice :	Oukid Lamia	MCB	Université Saad dahleb
Promotrice :	Madani Amina	MCB	Université Saad dahleb
Co-promotrice :	Boumahdi Fatima	MCB	Université Saad dahleb

Année : 2019-2020

Remerciements

Nous souhaitons d'abord remercier notre promotrice, Mme Madani Amina pour sa contribution dans la réalisation de ce travail, nous remercions Mme Boumahdi Fatima pour tout son soutien, sa patience et son aide pendant toute la durée de la recherche.

Que Mme Oukid Lamia, Maître de conférences à l'USDBlida1, trouve nos plus vifs remerciements pour l'honneur qu'elle nous fait en acceptant la présidence du jury.

A Mme Hireche Célia, Maître assistante à l'USDBlida1, nous la remercions d'avoir accepté d'examiner ce travail.

Nous souhaitons également témoigner notre gratitude envers Mme Boukenaoui Ferrouk Nouria et M Ferrouk Mostapha pour leur soutien qui nous a été précieux pour mener à bien ce projet de fin d'étude.

Nous remercions nos chers parents pour leur contribution, leur encouragement, et leur croyance.

Nous remercions nos sœurs, et merci à toutes nos familles.

Nous tenons à nous remercier individuellement pour les encouragements et l'épaullement dans les moments les plus difficiles.

Enfin, nous tenons à remercier tous ceux qui, de près ou de loin, ont contribué à la réalisation de ce travail.

Résumé

Plusieurs études dans la littérature ont montré que les mots dont les gens utilisent sont indicatifs de leurs états psychologiques. En particulier, des spécialistes en psychologie ont constaté que la dépression était associée à des modèles linguistiques distinctifs. La dépression est une maladie grave qui déstabilise l'organisme humain, elle peut conduire jusqu'au suicide. Une estimation précoce de la dépression peut aider à réduire et atténuer les risques sanitaires associés. Les réseaux sociaux comme Twitter, Reddit, Facebook, attirent les gens à faire part de leurs expériences et leurs opinions. Ces plateformes peuvent fournir des données massives, réelles et disponibles. Notre étude est orientée vers l'estimation du niveau de dépression à partir des publications Reddit en essayant de répondre sur le questionnaire de l'Inventaire de dépression de Beck (BDI). Nous avons développé trois modèles différents (CNN, Bi-LSTM et la combinaison des deux) issus de l'apprentissage profond avec des techniques statistiques infallibles pour estimer la sévérité de la dépression. Nos résultats concluons ont été évalué à la fin de ce mémoire.

Mots clé : Sévérité de dépression, Reddit, l'Inventaire de dépression de Beck, apprentissage profond, CNN, biLSTM, traitement du langage naturel.

Abstract

Several studies in the literature have shown that the words people use are indicative of their psychological states. In particular, psychologists have found that depression is associated with distinctive language patterns. Depression is a serious illness which destabilizes the human organism, it can lead to suicide. Early depression estimation can help to reduce and mitigate the associated health risks. Social networks like Twitter, Reddit, Facebook attract people to share their experiences and opinions. These platforms can provide massive and real and available data. Our study was directed towards estimating the level of depression from Reddit publications by attempting to answer the Beck Depression Inventory (BDI) questionnaire. We have developed three different models (CNN, Bi-LSTM and the combination of the two) derived from deep learning with infallible statistical techniques to estimate the severity of depression. Our final results were evaluated at the end of this thesis.

Keywords: Depression severity, Reddit, Beck's Depression Inventory, deep learning, CNN, biLSTM, natural language processing.

ملخص

أظهرت العديد من الدراسات في الأدبيات أن الكلمات التي يستخدمها الناس تدل على حالاتهم النفسية وجد علماء النفس على وجه الخصوص أن الاكتئاب مرتبط بأنماط لغوية مميزة. الاكتئاب مرض خطير يزعزع استقرار الكائن البشري ويمكن أن يؤدي إلى الانتحار. يمكن أن يساعد التقدير المبكر للاكتئاب في تقليل المخاطر الصحية المرتبطة بها والتخفيف من حدتها. تجذب الشبكات الاجتماعية مثل تويتر وريدت , فايس بوك الأشخاص لمشاركة تجاربهم وآرائهم. يمكن أن توفر هذه المنصات بيانات ضخمة حقيقية ومتاحة. تم توجيه دراستنا نحو تقدير مستوى الاكتئاب من مشاركة محاولة لإجابة لقد قمنا بتطوير ثلاثة نماذج مختلفة مستمدة *CNN, biLSTM, CNN_biLSTM* على استبيان من التعلم العميق بتقنيات إحصائية مضمونة لتقدير شدة الاكتئاب. تم تقييم نتائجنا النهائية في نهاية هذه الأطروحة.

الكلمات المفتاحية:

معالجة اللغة الطبيعية ، *biLSTM* ، *CNN* ، التعلم العميق ، Beck's Depression Inventory ، Reddit ، شدة الاكتئاب

Sommaire

Introduction générale.....	1
Chapitre 01 : Analyse des sentiments	4
1. Introduction.....	4
2. Définitions de l'analyse des sentiments.....	4
3. Catégories d'applications de l'analyse des sentiments et les domaines d'applications.	5
3.1. Détection.....	5
3.2. Prédiction	6
4. Types d'analyse des sentiments.....	6
4.1. Analyse fine des sentiments (Fine-grained Sentiment Analysis).....	6
4.2. Détection d'émotion (Emotion detection)	7
4.3. Analyse de sentiments à base d'aspects (Aspect-Based Sentiment Analysis ABSA)	7
5. Classification de l'analyse des sentiments	7
5.1. Approche automatique.....	7
5.2. Approche basée sur des regles (Rules-Based Approach)	8

Sommaire

5.3. Approche basée sur le lexique (Lexicon-Based Approach).....	8
5.4. Approche hybride	10
6. Niveaux de classification	10
6.1. Niveau du document (Message level ou Document level)	10
6.2. Niveau de la phrase (Sentence level)	10
6.3. Niveau des aspects (Feature level)	10
7. Avantages de l'analyse de sentiment.....	10
8. Traitement automatique du langage naturel.....	11
9. Conclusion	11
Chapitre 02 : Apprentissage en profondeur	13
1. Introduction.....	13
2. Apprentissage automatique (Machine Learning) ML	13
3. Méthodes d'apprentissage automatique.....	13
3.1. Apprentissage non supervisé (Unsupervised Learning)	14

Sommaire

3.2. Apprentissage par renforcement (Reinforcement Learning) :.....	15
4. Apprentissage par transfert (Transfere Learning) :.....	15
5. Réseaux de Neurones artificiels (Artificial Neural Network)	15
6. Apprentissage profond (Deep Learning) DL	16
6.1. Perceptron	17
6.2. Bias.....	17
6.3. Séparabilité linéaire et non linéaire	18
6.4. Fonction d'activation	19
6.6. Fonction de Perte (Loss Function)	21
6.7. Perceptron multicouche (Multilayer Perceptron) MLP	22
6.8. Réseaux de Neurones convolutifs (Convolutional neural networks) CNN	22
6.9. Réseaux neuronaux à action directe (Feed-Forward Neural Networks).....	23
6.10. Réseaux de neurones récurrents (<i>Recurrent Neural Networks</i>) RNN	23
6.11. Mémoire à long/court terme (Long Short Term Memory) LSTM	24

Sommaire

6.12. Mémoire à long terme à court terme bidirectionnelle (Bidirectional Long Short Term Memory) BILSTM (détaillé de l'architecture)	25
7. Conclusion	25
Chapitre 03 : La mesure du degré de dépression : état de l'art	27
1. Introduction.....	27
2. Troubles mentaux.....	27
3. Troubles dépressifs majeurs (mettre à jour de l'article)	28
4. Exploration des réseaux sociaux pour détecter les personnes dépressives	29
5. Mesure du degré de dépression sur les réseaux sociaux	31
6. Discussion	36
7. Conclusion	37
Chapitre 04 : Conception d'un modèle de mesure du degré de dépression à partir des réseaux sociaux.	39
1. Introduction.....	39

Sommaire

2. Architecture de notre modèle.....	39
3. Extraction de données (Training data, Test data)	41
4. Prétraitement et vectorisation.....	42
5. Apprentissage des modèles	46
6. Test (Prediction) :	51
7. Conclusion :	53
Chapitre 05 : Expérimentations et résultats	55
1. Introduction.....	55
2. Matériel.....	55
3. Logiciel.....	55
4. Librairies	56
5. Model Word2Vec :	57
6. Data Set :	57
7. Code source :	58

Sommaire

8. Reddit :	63
9. Mesure de performance	64
10. Résultat et évaluations des performances	65
11. Conclusion	71
Conclusion générale	73
Annexe	75
Références bibliographiques	83

Liste des figures

Figure 1 Etapes de l'approche automatique	8
Figure 2 Etapes de l'approche à base de règles	8
Figure 3 (a) Sans Bias.....	17
Figure 4 Avec Bias.....	18
Figure 5 Fonction softmax [51].	21
Figure 6 Différentes Couches du CNN	22
Figure 7 Architecture de notre modèle de mesure du niveau de la dépression	40
Figure 8 Structure des fichiers XML du Dataset de eRisk 2020.....	41
Figure 9 Structure d'un fichier d'apprentissage csv.....	41
Figure 10 Structure d'un fichier de test csv	42
Figure 11 Exemple d'une phrase avec Stopwords	43
Figure 12 Phrase de la figure 10 sans stopwords.....	43
Figure 13 Phrase de la figure 11 après lemmatisation.....	43

Figure 14 Tokens de la phrase de la figure 12.....	44
Figure 15 Exemple sur l'utilité du padding.....	44
Figure 16 Exemple d'une représentation vectorielle	45
Figure 17 Modèle Skip-gram et CBOW [1]	46
Figure 18 Un réseau neuronal convolutif pour le TAL	47
Figure 19 Architecture de notre modèle CNN	49
Figure 20 Architecture de notre modèle BiLSTM.....	50
Figure 21 Architecture de notre modèle BiLSTM_CNN	51
Figure 22 Code source 1.....	59
Figure 23 Code source 2.....	59
Figure 24 Code source 3.....	60
Figure 25 Code source 4.....	60
Figure 26 Code source 5.....	61
Figure 27 Code source 6.....	61

Figure 28 Code source 7.....	62
Figure 29 Code source 8.....	62
Figure 30 Code source 9.....	63
Figure 31 Code source 10.....	63

Liste des diagrammes

Diagramme 1 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure AHR au challenge eRisk2020	66
Diagramme 2 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure ACR au challenge eRisk2020	67
Diagramme 3 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure ADODL au challenge eRisk2020	67
Diagramme 4 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure DCHR au challenge eRisk2020	68

Liste des tableaux

Tableau 1 Avantages et des limites des deux approches précédentes basées sur différents niveaux d'analyse du sentiment [24-26]	10
Tableau 2 Comparaison entre le max pooling et l'average pooling.....	23
Tableau 3 Avantages et inconvénients de l'architectures de RNN	24
Tableau 4 Comparaison entre les travaux de mesure du niveau de dépression.....	36
Tableau 5 Exemple de prédiction de chaque publication d'une personne	52
Tableau 6 Structure du fichier des réponses du questionnaire de "BDI"	58
Tableau 7 montrant la performance des résultats eRisk2020	69
Tableau 8 récapitulant Nos Resultats eRisk2019	70
Tableau 9 La performance des résultats eRisk 2019 associés à nos résultats [5].	71

Introduction générale

Selon l'Organisation mondiale de la santé environ 450 millions de personnes dans le monde souffrent actuellement de problèmes liés aux troubles mentaux, ce qui place les troubles mentaux dans les causes principales de morbidité et d'incapacité à l'échelle mondiale[1]. La dépression compte parmi les **dix pathologies majeures du vingt et unième siècle**. Le manque de soutien aux personnes ayant des troubles mentaux, associés à la peur de la stigmatisation, empêche beaucoup de personnes d'accéder aux traitements dont elles ont besoin pour mener des vies productives et en bonne santé. Une détection précoce est un but qui peut améliorer le traitement et les résultats.

Pour y remédier face à ce problème des chercheurs ont abouti à une solution qui consiste à exploiter les plateformes de médias sociaux en ligne dû à leur massive quantité de texte qui peut être trouvé dedans et de leur disponibilité, car ses plateformes permettent aux gens de partager leurs pensées et dialoguer ouvertement avec les autres. L'une des hypothèses derrière de nombreuses tâches de calcul dans le domaine psychologique est que le texte en dit beaucoup de choses sur l'écrivain. Par conséquent, la prédiction des traits psychologiques des personnes à partir du texte est devenue un domaine de recherche important.

Ce travail vise à exploiter les messages échangés dans les réseaux sociaux afin de mesurer la gravité des signes de dépression à partir des réseaux sociaux. Compte tenu de l'historique des messages, l'objectif est de pouvoir estimer automatiquement la gravité des symptômes associés à la dépression. Nous utilisons les publications de Reddit de eRisk 2020 (Early risk prediction on the Internet) fournies par la conférence CLEF (Conference and Labs of the Evaluation Forum)¹.

Cette étude permet d'utiliser les techniques d'apprentissage en profondeur pour évaluer divers symptômes de la dépression en répondant sur 21 questions de l'inventaire de dépression de beck « BDI », sur la base des preuves trouvées dans les écrits des utilisateurs. Nous avons proposé trois approches basées sur des modèles d'apprentissage en profondeur pour remplir le questionnaire BDI. Pour la première fois, les réseaux de neurones sont utilisés pour mesurer la gravité des signes de dépression. Par le biais des modèles de réseaux de neurones convolutifs (CNN) et des modèles de mémoire bidirectionnelle à long court terme (BiLSTM) aussi par la combinaison de ses deux derniers. Par la suite, pour générer différentes exécutions, nous avons utilisés deux méthodes statistiques afin d'estimer la sévérité de dépression.

La structure de ce mémoire est organisée comme suit :

¹ <https://early.irlab.org/>

Introduction Générale

Chapitre 01 : présente l'analyse des sentiments, ses domaines d'application et domaines associés, ses différents types et enfin certaines de ses techniques les plus utilisées.

Chapitre 02 : nous commençons par présenter le deep learning et ses notions de base, puis nous passons aux différentes techniques de deep learning utilisées dans l'analyse des sentiments.

Chapitre 03 : Le troisième chapitre est consacré aux travaux liés à la détection de la dépression qui abordent la même problématique ou une problématique similaire à la nôtre.

Chapitre 04 : Dans le quatrième chapitre, nous présentons notre modèle proposé et nous explorons chacune de ses couches tout en expliquant ses origines et les différents modèles sur lesquels il est construit

Chapitre 05 : Dans le dernier chapitre, nous décrivons le processus qu'on a suivie pour former et évaluer notre modèle étape par étape . À la fin, nous partageons les résultats que nous avons obtenus et nous les comparons à d'autres travaux.

Chapitre 01

Analyse des sentiments

Chapitre 01 : Analyse des sentiments

1. Introduction

Initialement l'analyse de sentiment a été introduite par Nasukawa et Yi[2], qui est l'un des sujets les plus brûlants de la recherche depuis la dernière décennie [3, 4], depuis, ce domaine est en accroissement exponentiel, car les opinions sont des éléments clés qui influencent au quotidien nos comportements, nos croyances et nos perceptions de la réalité. Les choix que nous faisons sont dans une large mesure conditionnée par la façon dont les autres voient et évaluent le monde, c'est pour cette raison que lorsque nous devons prendre une décision nous recherchons souvent les opinions des autres.

L'évolution de l'analyse des sentiments est directement liée à l'avancement des médias sociaux comme les forums, les blogs, les réseaux sociaux comme Twitter, Reddit, Facebook. Ces médias attirent les gens à faire part de leurs expériences et leurs opinions, quelques trillions de messages, de photos et des messages vocaux sont envoyés dans le monde entier chaque jour.

Dans ce chapitre, nous introduisons l'analyse des sentiments ainsi que ses domaines d'application et ses différents niveaux, à la fin nous nous intéressons dans ce chapitre aux méthodes multidimensionnelles de l'analyse d'opinion en particulier le traitement automatique de la langue.

2. Définitions de l'analyse des sentiments

L'analyse des sentiments définit des outils automatisés qui extraient des informations subjectives de textes constituant des opinions et des sentiments de manière à créer des connaissances structurées pouvant être facilement utilisées pour la prise de décision.

Plusieurs synonymes existent tels que l'exploration de l'opinion, fouille d'opinion ou même opinion mining en anglais. Le terme principal qui se répète est l'opinion. D'après Larousse l'opinion : est un jugement, avis, sentiment qu'un individu ou un groupe émet sur un sujet, des faits, ce qu'il en pense. L'opinion peut être définie comme l'expression des sentiments d'une personne envers une entité. Elle est subjective et peut être décrite avec certains attributs. L'attribut d'opinion le plus étudié est la polarité qui est soit positive ou négative et peut être éventuellement neutre, c'est ce qui définit si l'opinion est favorable ou défavorable. D'autres attributs sont l'intensité de l'opinion et le degré de subjectivité.

L'analyse de sentiment analyse les points de vue, les sentiments, les évaluations, le comportement et la psychologie des personnes à l'égard des entités vivantes et abstraites, le sentiment fait référence à l'attitude mentale, qui est créée sur la base de l'émotion. Ce sentiment est utilisé pour transmettre les pensées de l'individu dérivant de son émotion. Contrairement à l'émotion, le sentiment va plus loin, il ne se limite pas aux dimensions

psychologiques. McDougall a déclaré que le sentiment relie généralement l'émotion primaire à l'action, l'analyse d'opinions est un domaine de recherche qui se concentre sur l'identification et la classification des opinions dans les données textuelles. L'analyse des sentiments peut cependant également reposer sur d'autres éléments que les données textuelles. Elle peut par exemple être basée sur l'usage des émoticônes (emojis), sur les « émotions », sur l'analyse de la voix ou même sur le facial coding / decoding . Elle est très étroitement liée à la recherche sur la linguistique informatique pour l'identification et l'extraction systématique d'informations thématiques, et le traitement des langues naturelles (NLP)[5].

3. Catégories d'applications de l'analyse des sentiments et les domaines d'applications

Avec l'émergence du web participatif (Web 2.0), tous les utilisateurs d'Internet ont la possibilité de donner leurs avis sur une infinité de sujets (produits, services, prestations d'entreprise) et d'en discuter sur des forums, blogs, ou des groupes de discussion. Toutes ces opinions sont cruciales pour les entreprises, tout d'abord afin de développer des tâches de veille (technologique, marketing, concurrentielle, sociétale), mais aussi pour détecter les consommateurs mécontents, déterminer la perception des gens sur de nouveaux produits ou se renseigner sur la réputation de leur société. Toutes ces données peuvent aussi servir aux acteurs du monde politique afin de connaître les intentions des citoyens.

L'analyse des sentiments a son importance et y est appliquée dans de nombreux domaines, la plupart des applications d'analyse des sentiments appartiennent aux deux catégories suivantes : la détection et la prédiction.

3.1. Détection

Elle consiste à étudier certains phénomènes ou de comprendre quelque chose qui s'est passé dans le monde et de trouver les raisons derrière elle, c'est pourquoi elles se concentrent principalement sur le passé et le présent des événements, en utilisant plusieurs méthodes de classification pour le but de déterminer la polarité des documents pertinents et les réordonner selon un score d'opinion.

Nous citons quelques exemples qui sont liés à cette catégorie :

Technologique : Innovation des logiciels, les développeurs prennent en considération les avis des internautes, ils ont un besoin constant de détecter les consommateurs mécontents ou pour déterminer la perception des gens sur de nouveaux produits [6].

Médiatique : Le flux de sentiments dans les réseaux sociaux a été étudié [7].

Analyse des sentiments

Sportif : Les relations entre la ligne de pari de la NFL (The National Football League) avec l'opinion publique dans les blogs et Twitter ont été étudiées [8].

Commerciale : Les revues ont été utilisées pour classer les produits et les marchands [9].

Politique : Le sentiment dans Twitter était lié aux sondages d'opinion publique [10].

3.2. Prédiction

Consiste à estimer une valeur, elle utilise la détection pour prédire les résultats futurs, la prédiction est très importante dans le monde des affaires et politique, et même dans le domaine de la santé.

Nous citons quelques exemples qui sont liés à cette catégorie :

Politique : Le sentiment Facebook a également été appliqué pour prédire les résultats des élections [11].

Sociale : Dans [12], l'analyse des sentiments a été utilisée pour caractériser les relations sociales.

Économique : Dans [13], un modèle de sentiment a été proposé pour prédire le rendement des ventes. Les données de Twitter, les critiques de films et les blogues ont été utilisés pour prédire les recettes au guichet des films[14-16].

Industriel : Les humeurs de Twitter ont été utilisées pour prédire le marché boursier [17].

Éducation : L'analyse des sentiments peut être utilisée pour extraire des informations utiles sur la méthodologie d'enseignement d'un enseignant et également sur le programme du cours. Il identifie le degré d'apprentissage des étudiants, comprend leurs besoins, prévoit leurs performances et apporte des changements effectifs dans le style. Les résultats de l'analyse des sentiments aident les enseignants et les établissements à prendre des mesures correctives [18].

4. Types d'analyse des sentiments

Il existe plusieurs types allant des systèmes qui se concentrent sur la classification de la polarité aux systèmes qui détectent des émotions ou identifient des intentions.

Nous allons citer les types d'analyse de sentiments les plus populaires :

4.1. Analyse fine des sentiments (Fine-grained Sentiment Analysis)

Consiste à classer les sentiments généraux d'un texte selon 5 catégories (très positive, positive, neutre, négative, très négative).

Analyse des sentiments

Ceci est généralement appelé analyse fine des sentiments et peut être utilisé pour interpréter les notes 5 étoiles dans un avis, par exemple: on attribue 5 étoiles si c'est très positif, et une étoile si c'est très négatif comme avis.

4.2. Détection d'émotion (Emotion detection)

La détection des émotions vise à détecter des émotions telles que le bonheur, la frustration, la colère, la tristesse, etc. De nombreux systèmes de détection d'émotions sont basés sur l'utilisation de lexiques de sentiments (c'est-à-dire des listes des émotions) ou sur des algorithmes d'apprentissage automatique complexes.

4.3. Analyse de sentiments à base d'aspects (Aspect-Based Sentiment Analysis ABSA)

Au lieu de classer le sentiment général d'un texte en positif ou en négatif, l'analyse de sentiments à base d'aspects permet d'analyser le texte afin d'identifier différents aspects et de déterminer le sentiment correspondant pour chacun. Les résultats sont plus détaillés, intéressants et précis, car l'analyse à base d'aspects examine de manière précise les informations contenues dans un texte.

5. Classification de l'analyse des sentiments

L'analyse des sentiments utilise divers méthodes et algorithmes de traitement du langage naturel (TAL), que nous détaillerons plus en détail dans cette section.

Les principaux types d'algorithmes utilisés comprennent :

- **Systèmes automatiques** : ils s'appuient sur des techniques d'apprentissage automatique pour apprendre à partir des données.
- **Systèmes basés sur des règles** : ils effectuent une analyse des sentiments basée sur un ensemble de règles créées manuellement.
- **Systèmes basés sur le lexique** : ils effectuent une analyse des sentiments basée sur le lexique.
- **Des systèmes hybrides** : ils combinent à la fois des approches basées sur des règles et des approches automatiques.

5.1. Approche automatique

Systèmes qui s'appuient sur des techniques d'apprentissage automatique à partir de données.

Analyse des sentiments

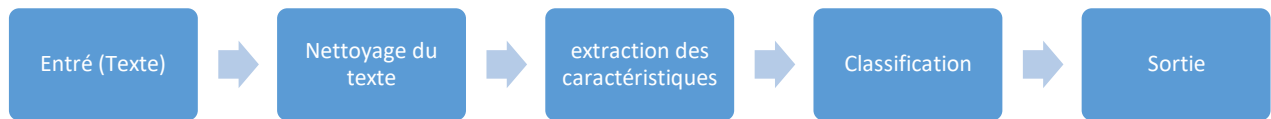


Figure 1 Etapes de l'approche automatique

5.2. Approche basée sur des règles (Rules-Based Approach)

L'approche basée sur des règles est utilisée en définissant diverses règles pour obtenir l'opinion, créées en marquant chaque phrase dans chaque document, puis en testant chaque jeton ou mot pour sa présence. Si le mot est là et a un sentiment positif, une note +1 lui a été appliquée. Chaque publication commence par un score neutre de zéro et a été considérée comme positive. Si le score de polarité final était supérieur à zéro, ou négatif si le score global était inférieur à zéro. Après la sortie de l'approche basée sur des règles, il vérifiera ou demandera si la sortie est correcte ou non [19]. Si la phrase d'entrée contient un mot qui n'est pas présent dans la base de données et qui peut aider à l'analyse de la critique de film, alors ces mots doivent être ajoutés à la base de données. Il s'agit d'un apprentissage supervisé dans lequel le système est formé pour apprendre si une entrée est donnée.

5.3. Approche basée sur le lexique (Lexicon-Based Approach)

L'analyse des sentiments basée sur le lexique fait référence à l'étude menée par les experts linguistiques. Le résultat de cette étude est un ensemble de lexique selon lesquelles les mots classés sont soit positifs, soit négatifs avec leur mesure d'intensité correspondante.

Voici un exemple de base du fonctionnement d'un système basé sur le lexique:

1. On a deux listes de mots polarisés (par exemple, des mots négatifs tels que mauvais, pire, laid, etc., et des mots positifs tels que bon, meilleur, beau, etc.).
2. On compte le nombre de mots positifs et négatifs qui apparaissent dans un texte donné.
3. Si le nombre d'apparitions de mots positifs est supérieur au nombre d'apparitions de mots négatifs, le système renvoie un sentiment positif et vice versa. Si les nombres sont pairs, le système renverra un sentiment neutre.

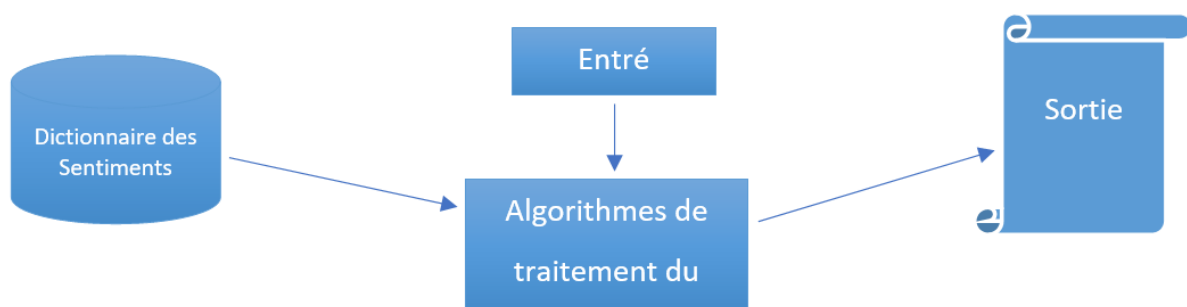


Figure 2 Etapes de l'approche à base de règles

Analyse des sentiments

Les systèmes basés sur des règles sont très naïfs, car ils ne prennent pas en compte la manière dont les mots sont combinés dans une séquence. Bien entendu, des techniques de traitement plus avancées peuvent être utilisées et de nouvelles règles peuvent être ajoutées pour prendre en charge de nouvelles expressions et un nouveau vocabulaire. Cependant, l'ajout de nouvelles règles peut affecter les résultats précédents et l'ensemble du système peut devenir très complexe. L'approche basée sur le lexique peut être divisée en deux catégories :

5.3.1. Basé sur un corpus d'apprentissage

Les approches basées sur des corpus reposent sur l'orientation contextuelle des mots d'opinion [20]. Dans cette technique, un corpus de mots d'opinion d'un domaine spécifique est créé à partir d'un groupe de mots-clés de base. La taille augmente en ajoutant d'autres terminologies d'opinion à l'aide de différentes techniques statistiques ou sémantiques.

5.3.2. Basé sur un dictionnaire

Elle s'appuie sur l'extraction de l'opinion puis l'utilisation du dictionnaire des synonymes et des antonymes. Cette technique basée sur l'orientation sémantique des mots, pour cela elle utilise un dictionnaire préconstruit tel que SentiWordNet qui est l'un des dictionnaires standards les plus utilisés de nos jours[21-23].

Le tableau suivant montre les avantages et les limites de la méthode basée sur l'apprentissage automatique et de l'approche basée sur un lexique :

Méthode d'analyse	Par apprentissage automatique	Lexicale
Avantages	1. Pourrait être transformé en ce que le domaine demande pour mieux travailler. 2. Un dictionnaire n'est pas nécessaire. 3. Donne de meilleurs résultats en termes de haute précision de classification.	1. Ne demande aucune donnée d'entraînement ou des données étiquetées et ceci permet d'introduire moins d'opérations de calcul
Inconvénients	1. Elle peut être affectée par les variations de classes et aussi par l'effet des changements linguistiques. 2. Les classificateurs qui se sont entraînés sur un domaine spécifique dans la plupart des cas ne fonctionnent pas avec un autre.	1. Moins de capacité de classification en fonction du contexte ou du domaine. 2. Exigeait l'existence de ressources linguistiques puissantes qui ne sont pas toujours disponibles.

Tableau 1 Avantages et des limites des deux approches précédentes basées sur différents niveaux d'analyse du sentiment [24-26]

5.4. Approche hybride

Systèmes combinant à la fois des approches basées sur des règles et des approches automatiques. Un énorme avantage de ces systèmes est que les résultats sont souvent plus précis.

6. Niveaux de classification

L'analyse des sentiments peut également être classée en fonction de la manière dont les opinions à analyser sont identifiées[27]. Les trois principaux niveaux de classification sont :

6.1. Niveau du document (Message level ou Document level)

Détermine la polarité d'un texte entier. Cette analyse fonctionne bien pour des documents qui présentent un point de vue précis, mais moins pour des comparaisons, car elle ne fera pas la différence entre les sujets abordés. [28-30].

6.2. Niveau de la phrase (Sentence level)

Détermine la polarité de chaque phrase contenue dans un texte, cette analyse peut donner une mesure de la "neutralité" d'un texte par exemple pour analyser des entrées de Wikipédia. Les méthodes utilisées sont celles de l'analyse de subjectivité puisqu'une phrase peut être subjective ou objective. La phrase objective contient les faits. Il n'a pas de jugement ou d'opinion sur l'objet ou l'entité alors que la phrase subjective a des opinions [31].

6.3. Niveau des aspects (Feature level)

Effectue une analyse plus fine que les autres niveaux, au lieu de déterminer les entités à analyser en fonction de critère structurale (phrase, paragraphe, document) ces méthodes se basent sur l'analyse de corrélation entre l'opinion émise et la cible de cette opinion. Par exemple, la phrase "Le sujet du cours me passionne, mais le professeur est ennuyeux." Présente deux sentiments sur l'entité "cours" : le sujet qui est perçu comme positif et le professeur, qui est perçu comme négatif. Ce niveau d'analyse permet de différencier les aspects qui sont aimés ou non par les auteurs des textes et ainsi permet plus facilement de déterminer des remédiations possibles. En revanche il est très difficile à mettre en place, car ce type d'analyse est extrêmement complexe [32].

7. Avantages de l'analyse de sentiment

- **Tri des données à grande échelle :** L'analyse des sentiments aide les entreprises à traiter d'énormes quantités de données de manière efficace et rentable.

- **Analyse en temps réel** : L'analyse des sentiments peut identifier les problèmes critiques en temps réel, par exemple un client en colère qui est sur le point de se désister, Les modèles d'analyse des sentiments peuvent nous aider à identifier immédiatement ces types de situations, afin qu'on puisse agir immédiatement.

8. Traitement automatique du langage naturel

L'analyse des sentiments se base essentiellement sur la classification des sentiments en trois catégories (positif, négatif, et neutre) selon [33] il existe deux approches principales :

- Approche basée sur l'apprentissage automatique
- Approche basée sur le TALN

Les techniques de l'apprentissage automatique seront détaillées dans le prochain chapitre, dans ce qui suit on va détailler l'approche basée TALN.

Le traitement automatique du langage naturel est basé sur le traitement des phrases, tandis que les phrases sont constituées de mots. Dans l'analyse de sentiment, on trouve des mots qui sont utilisés pour exprimer l'opinion, un lexique ou vocabulaire bien déterminé est employé dans la classification de sentiment. Il y a des mots d'opinion positifs qui sont utilisés pour exprimer des sentiments positifs tel que : il me plaît, beau, joli, intéressant ... et des mots d'opinion négatifs qui sont utilisés pour exprimer les sentiments négatifs. Il y a aussi des phrases d'opinion et des idiomes qui ensemble sont appelés lexique d'opinion. Il y a trois principales approches afin de compiler ou de recueillir la liste des mots d'opinion. L'approche manuelle prend beaucoup de temps et n'est pas utilisée seule. Elle est généralement combinée avec les deux autres automatisés approches comme une vérification finale pour éviter les erreurs qui ont résulté à partir de méthodes automatisées [34].

9. Conclusion

Dans ce chapitre nous avons fait un survol sur les principaux fondements, méthodes et techniques de classification des sentiments et d'opinion mining et la relation avec le traitement automatique du langage naturel. Dans le chapitre suivant, nous allons détailler l'apprentissage profond.

Chapitre 02

Apprentissage en profondeur

Chapitre 02 : Apprentissage en profondeur

1. Introduction

Ces dernières années, l'évolution de l'apprentissage automatique a conduit à des progrès importants et des améliorations dans la façon dont nous interagissons avec notre monde. L'une de ces captivantes progressions est le domaine de l'apprentissage en profondeur (Deep Learning). Dans ce chapitre, nous présenterons ce domaine fastidieux, nous débiterons par la définition des termes les plus importants et les plus utilisés tout au long de ce mémoire afin de faciliter la compréhension des différentes parties de notre modèle proposé. Ensuite nous aborderons les différentes classes et techniques de l'apprentissage en profondeur en nous appuyant sur des exemples.

2. Apprentissage automatique (Machine Learning) ML

L'apprentissage automatique « ML » est un sous-domaine de l'intelligence artificielle (IA). Il englobe la conception, l'analyse, le développement et l'implémentation des méthodes permettant à une machine de progresser par un processus systématique, afin de remplir des tâches difficiles ou problématiques par des moyens algorithmiques. Il adapte ses analyses et ses comportements en solution définie, en se basant sur l'analyse de données empiriques² provenant d'une base de données ou de capteurs [35].

3. Méthodes d'apprentissage automatique

On distingue quelques méthodes d'apprentissage automatiques : apprentissage supervisé, apprentissage non supervisé et l'apprentissage par renforcement qui sera détaillé par la suite.

3.1. Apprentissage supervisé (Supervised Learning)

« SL » Référence à un travail avec des classes prédéterminées et des exemples connus, d'où le système apprend à classer à partir d'un modèle connu on parle d'analyse discriminante (apprentissage supervisé).

Exemple : si nous essayons d'apprendre le sentiment ou l'impression d'un film, l'ensemble de données peut être un ensemble de critiques de films et les étiquettes sont de 0 à 5 étoiles.

Il existe deux types d'apprentissages supervisés qui sont :

² C'est la recherche basée sur l'expérimentation dont l'objectif est de tester une hypothèse.

3.1.1. Classification

Consiste à mapper une entrée dans un ensemble fixe de catégories, autrement dit les sorties sont des catégories d'étiquette.

Exemple : la classification d'une image de chat ou de chien.

Les problèmes de classification peuvent être divisés en deux catégories :

- Classification binaire (Binary Classification) :

C'est la forme la plus simple de classification où le travail consiste à classer une entité dans l'une des deux catégories possibles.

Exemple : En donnant les attributs des fruits comme le poids, la couleur, la texture de la peau, etc. qui caractérise les fruits comme pêche ou pomme.

- Classification multi classes (Multi_Class) :

Une tâche de classification avec plus de deux classes. L'idée consiste à ce que chaque input il sera affecté à une et une seule étiquette.

Exemple : pour classer un ensemble d'images de fruits qui peuvent être des oranges, des pommes ou des poires. La classification multiclass fait l'hypothèse que chaque échantillon est affecté à une et une seule étiquette : un fruit peut être soit une pomme, soit une poire, mais pas les deux à la fois.

Tous les problèmes de classification peuvent être résolus avec le réseau de neurones artificiels que nous aborderons un peu plus en bas.

3.1.2. Régression :

Mappent une entrée à une valeur numérique réelle, en outre les sorties sont des nombres continus.

Exemple : essayer de prédire le coût d'une facture d'électricité ou du prix du stock du supermarché.

3.1. Apprentissage non supervisé (Unsupervised Learning)

L'apprentissage non supervisé « NSL » détermine les catégories à partir des données où il n'y a pas d'étiquettes présentes. Ces tâches peuvent prendre la forme de regroupements, regrouper des éléments similaires, ou similitudes, définissant le lien entre une paire d'éléments.

Apprentissage en profondeur

Exemple : Si nous voulions recommander un film en fonction des habitudes de visionnage d'une personne, nous pouvions utiliser le cluster en fonction de ce qu'ils ont regardé et apprécié, pour évaluer les habitudes de visionnement correspond le mieux à la personne à qui nous recommandons le film.

3.2. Apprentissage par renforcement (Reinforcement Learning) :

L'apprentissage par renforcement se concentre sur la maximisation d'une récompense donnée par une action ou un ensemble d'actions prises. Les algorithmes sont entraînés pour encourager certains comportements et décourager autres.

Exemple : L'apprentissage par renforcement a tendance à bien fonctionner sur des jeux comme les échecs ou le célèbre jeu GO, où la récompense peut faire gagner le match. Dans ce cas, un certain nombre d'actions doivent être prises avant que la récompense soit atteinte [36].

4. Apprentissage par transfert (Transfere Learning) :

L'idée de base est d'aider le modèle à s'adapter aux situations qu'il n'a pas rencontrées auparavant. Cette forme d'apprentissage repose sur la transformation d'un modèle vers un nouveau domaine. Apprendre de nombreuses tâches pour améliorer conjointement les performances à travers toutes les tâches c'est ce qui est appelé apprentissage multitâches. Ces techniques sont devenues propres à la fois sur l'apprentissage en profondeur et le TAL [36].

5. Réseaux de Neurones artificiels (Artificial Neural Network)

ANN est un modèle du ML dont le fonctionnement est inspiré des neurones biologiques et qui s'appuie également sur les méthodes statistiques. ANN est un modèle paramétrique qui a un ensemble de paramètres qui sont tirés du problème, tel que les poids et les biais. Il dispose également d'un certain nombre d'hyperparamètres qui peuvent être réglés par l'ingénieur, tels que le taux d'apprentissage et le nombre de couches cachées, il peut être créé pour produire des sorties continues ou discrètes, et il peut fonctionner avec des modèles linéaire ou non-linéaire aussi. D'après [37], Les ANN sont en fait des modèles linéaires qui sont regroupés pour résoudre des problèmes complexes. Il existe plusieurs types de réseaux neuronaux, dont chacun vient avec ses propres cas d'utilisation spécifiques et niveaux de complexité, nous citons ci-dessous deux types :

- Réseau neuronal FeedForward (RNFF): l'information voyage dans une seule direction de l'entrée vers la sortie.
- Réseau neuronal récurrent (RNN) : dans lequel les données peuvent circuler dans plusieurs directions.

Apprentissage en profondeur

Ces réseaux neuronaux possèdent de plus grandes capacités d'apprentissage et sont largement utilisés pour des tâches plus complexes comme l'apprentissage de l'écriture ou de la reconnaissance linguistique [38, 39].

Le bon choix du réseau de neurones dépend des données, et la spécificité de la tâche demandée.

6. Apprentissage profond (Deep Learning) DL

Le terme «apprentissage en profondeur» est quelque peu ambigu. C'est un terme de re-branding pour les réseaux de neurones ou est utilisé pour désigner les réseaux de neurones avec plusieurs couches consécutives (profondes).

Cependant le deep learning est un sous-domaine du machine learning. Profond est un terme technique. Il fait référence au nombre de couches dans un réseau neuronal. Un réseau peu profond a une couche cachée et un réseau profond en a plus d'une. Son fonctionnement est associé à celui de ANN, plusieurs couches cachées permettent aux réseaux de neurones profonds d'apprendre les caractéristiques des données dans une hiérarchie des fonctionnalités, car les entités simples telles que deux pixels se recombinent d'une couche à l'autre, pour former des entités plus complexes par exemple une ligne à travers des modules composés, simples, mais non linéaires. Les réseaux avec plusieurs couches transmettent les données d'entrée (caractéristiques) à travers des opérations mathématiques et sont donc plus exigeants en termes de calcul pour s'entraîner. L'intensité de calcul est l'une des caractéristiques de l'apprentissage profond [40].

Il existe de nombreuses applications réussies de l'apprentissage en profondeur dans de nombreux domaines, notamment la reconnaissance vocale, la traduction, le traitement d'images, la robotique, l'ingénierie, la science des données, la santé et bien d'autres. Nous citons quelque application populaire :

- Les chatbots qui prédisent la réponse en langage naturel sont disponibles depuis de nombreuses années. Les réseaux d'apprentissage en profondeur peuvent améliorer considérablement les performances des chatbots [41]. Aujourd'hui, ils fournissent des systèmes d'aide pour les entreprises et les aides à domicile comme Alexa d'Amazon et Google home.
- Google WaveNet (développé par DeepMind [42]), pour générer de la parole à partir du texte.
- Google Maps a été amélioré après le développement de l'apprentissage en profondeur pour analyser plus de 80 millions d'images Street View et extraire les noms des routes et des entreprises [43].
- Le diagnostic des soins de santé a été développé à l'aide du modèle Adversarial Auto-encoder qui a trouvé de nouvelles molécules pour lutter contre le cancer. L'identification et la génération de nouveaux composés sont basées sur les données biochimiques disponibles [44].

- Dans l'ingénierie plus traditionnelle, les applications scientifiques, comme l'analyse spatio-temporelle et financière, l'apprentissage en profondeur a montré des performances supérieures à celles des techniques d'apprentissage statistiques traditionnelles [45-50].

6.1. Perceptron

Le deep learning dans sa forme la plus simple est une évolution de l'algorithme perceptron formé avec un optimiseur basé sur un gradient. L'algorithme perceptron est l'un des premiers algorithmes d'apprentissage supervisé, datant des années 50 [51]. Tout comme un neurone biologique, l'algorithme perceptron agit comme un neurone artificiel, ayant plusieurs entrées et des poids associés à chaque entrée, dont chacune donne ensuite une sortie.

6.2. Bias

L'algorithme perceptron apprend un hyperplan qui sépare deux classes. Cependant, à ce stade, l'hyperplan de séparation ne peut pas s'éloigner de l'origine, bien que les données soient linéairement séparables [51] l'algorithme perceptron n'est pas en mesure de séparer les données. Cela est dû à la restriction du plan de séparation devant passer par l'origine exemple :

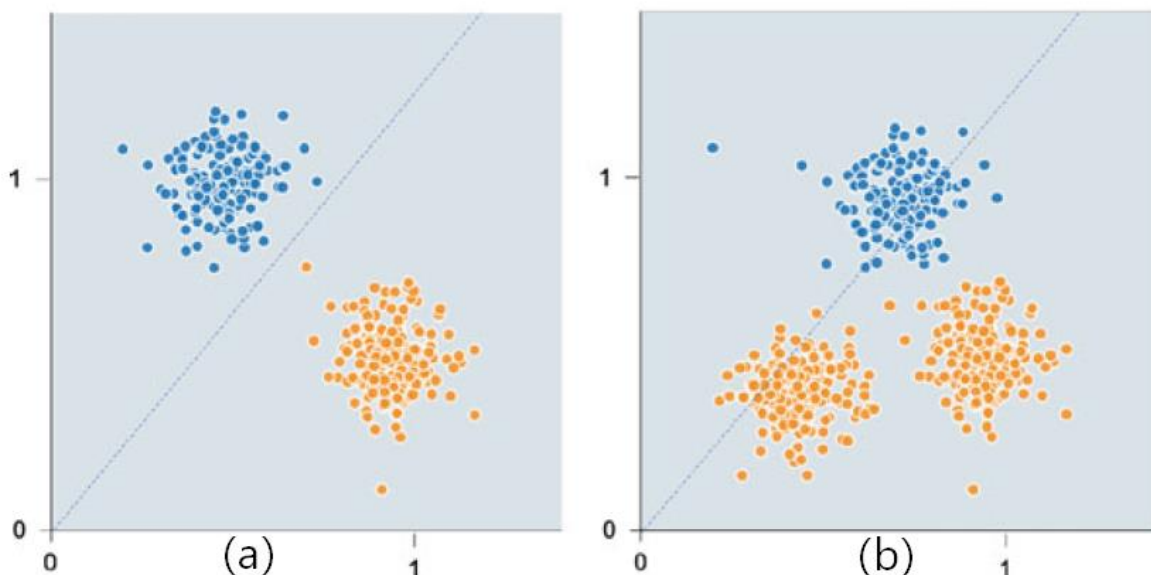


Figure 3 (a) Sans Bias

Une solution consiste à garantir que nos données peuvent être apprises si nous normalisons la méthode pour centrer autour de l'origine comme une solution potentielle pour atténuer ce problème ou ajouter un bias, permettant à l'hyperplan de classification de s'éloigner de l'origine.

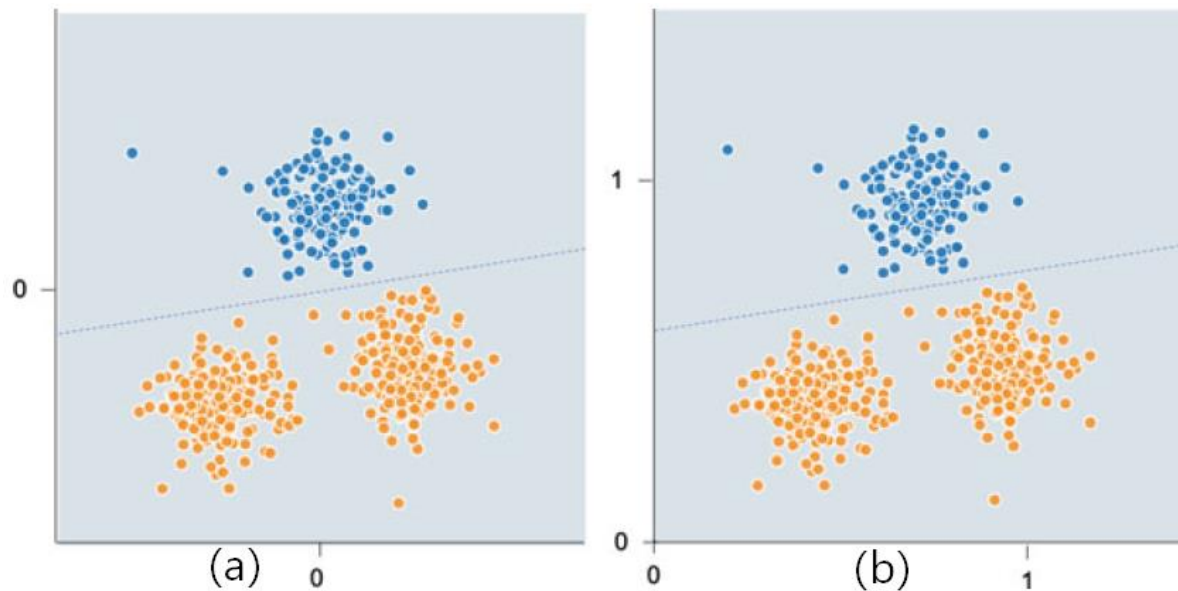


Figure 4 Avec Bias

Alternativement, nous pouvons traiter b comme un poids supplémentaire w_0 lié à une entrée constante de 1 :

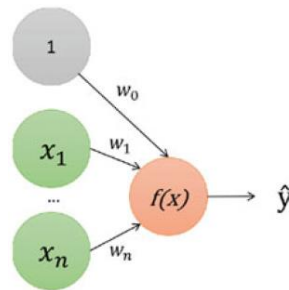


Figure 5 Diagramme de Perceptron incluant le biais

6.3. Séparabilité linéaire et non linéaire

Deux ensembles de données sont linéairement séparables si une seule frontière de décision peut les séparer.

Exemple : deux ensembles, A et B, sont linéairement séparables si pour un seuil de décision t , chaque $x_i \in A$ satisfait l'inégalité $\sum_i \omega_i x_i \geq t$ et chaque $y_i \in B$ satisfait $\sum_i \omega_i y_i < t$ avec ω_i égale aux poids

Inversement, deux ensembles ne sont pas séparables linéairement si la séparation nécessite une frontière de décision non linéaire.

Malheureusement, la plupart des données que nous avons tendance à rencontrer dans TAL et la parole sont très non linéaire. Pour ce faire il suffit de créer des combinaisons non linéaires des données d'entrée et les utiliser comme fonctionnalités dans le modèle. Une autre astuce consiste d'apprendre les fonctions non linéaires des données brutes, ce qui est l'objectif principal des réseaux de neurones [51].

6.4. Fonction d'activation

- Fonction d'activation linéaire : C'est une fonction simple de la forme: $f(x) = ax$ ou $f(x) = x$. En gros, l'entrée passe à la sortie sans une très grande modification ou alors sans aucune modification.
- Fonction d'activation non linéaire : Les fonctions non linéaires permettent de séparer les données non linéairement séparables. Elles constituent les fonctions d'activation les plus utilisées. Une équation non linéaire régit la correspondance entre les entrées et la sortie.

L'utilisation de fonctions d'activation non linéaires est tout simplement indispensable pour la simple et bonne raison que les fonctions linéaires ne fonctionnent qu'avec une seule couche de neurone [51]. Car au-delà d'une couche de neurones, l'application récurrente d'une même fonction d'activation linéaire n'aura plus aucun impact sur le résultat. Autrement dit, afin de résoudre des problèmes complexes, l'utilisation de fonctions non linéaires est indispensable.

6.5. Types de fonction d'activation

Étudiant maintenant les fonctions d'activation les plus utilisées [51] :

- Sigmoid : La fonction sigmoïde est une activation utile pour diverses raisons. Car elle réduit la valeur entre 0 et 1. En plus d'exprimer la valeur sous forme de probabilité, si la valeur en entrée est un très grand nombre positif, la fonction convertira cette valeur en une probabilité de 1. À l'inverse, si la valeur en entrée est un très grand nombre négatif, la fonction convertira cette valeur en une probabilité de 0. D'autre part, l'équation de la courbe est telle que, seules les petites valeurs influent réellement sur la variation des valeurs en sortie.

$$\text{Elle est définie comme : } \sigma(x) = \frac{1}{1+e^{-x}}$$

La fonction sigmoïde a plusieurs inconvénients :

- Elle n'est pas centrée sur zéro, c'est-à-dire que des entrées négatives peuvent engendrer des sorties positives.
- Étant assez plate, elle influe assez faiblement sur les neurones par rapport à d'autres fonctions d'activations. Le résultat est souvent très proche de 0 ou de 1 causant la saturation de certains neurones.

Apprentissage en profondeur

- Elle est coûteuse en termes de calcul, car elle comprend la fonction exponentielle.
- Tangente hyperbolique (Tanh) : Cette fonction ressemble à la fonction sigmoïde. La différence avec la fonction sigmoïde est que la fonction Tanh produit un résultat compris entre -1 et 1. La fonction Tanh est en terme général préférable à la fonction sigmoïde, car elle est centrée sur zéro. Les grandes entrées négatives tendent vers -1 et les grandes entrées positives tendent vers 1.

Elle est définie comme : $f(x) = \tanh(x)$

Mis à part l'avantage cité au-dessus, elle contient les mêmes inconvénients que la fonction sigmoïde.

- Unité de Rectification linéaire (ReLU) : Pour résoudre le problème de saturation des deux fonctions précédentes (Sigmoïde et Tanh) il existe la fonction ReLU. Cette fonction est la plus utilisée.

Elle est définie comme : $f(x) = \max(0, x)$

Si l'entrée est négative, la sortie est 0 et si elle est positive alors la sortie est x. Cette fonction d'activation augmente considérablement la convergence du réseau et ne sature pas.

Mais la fonction ReLU n'est pas parfaite. Si la valeur d'entrée est négative, le neurone reste inactif, ainsi les poids ne sont pas mis à jour et le réseau n'apprend pas.

- Softmax ; la régression softmax (ou régression logistique multinomiale) est une généralisation de la régression logistique au cas où nous voulons traiter plusieurs classes. Il est particulièrement utile pour les réseaux de neurones où l'on souhaite appliquer une classification non binaire. Dans ce cas, la simple régression logistique ne suffit pas. Nous aurons besoin d'une distribution de probabilités pour toutes les étiquettes, ce que Softmax nous donne.

Elle est définie comme : $\sigma(x_j) = \frac{e^{x_j}}{\sum_i e^{x_j}}$

On peut utiliser softmax pour produire un vecteur de probabilités en fonction de la sortie de ce neurone. Dans le cas d'un problème de classification qui a $K = 3$ classes, la couche finale de notre réseau sera une couche entièrement connectée avec une sortie de trois neurones. Si nous appliquons la fonction softmax à la sortie de la dernière couche, nous obtenons une probabilité pour chaque classe en attribuant une classe à chaque neurone.

Les probabilités softmax peuvent devenir très petites, surtout lorsqu'il y a de nombreuses classes et les prédictions deviennent plus confiantes. La fonction softmax est un cas particulier des fonctions d'activation, en ce qu'il est rarement considéré comme une activation qui se produit entre les couches. Par conséquent,

Apprentissage en profondeur

softmax est souvent traitée comme la dernière couche d'un réseau pour la classification multiclassées plutôt qu'une fonction d'activation.

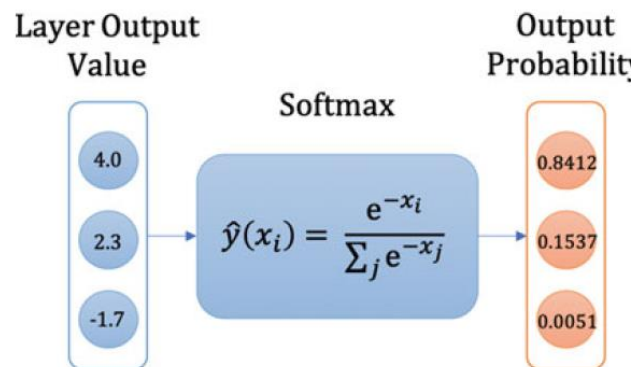


Figure 5 Fonction softmax [51].

Bien qu'il existe d'autres fonctions d'activation, nous avons cité les plus utilisées.

6.6. Fonction de Perte (Loss Function)

Un autre aspect important de la formation des réseaux de neurones est le choix des fonctions de perte. La sélection de la fonction d'erreur dépend du type de problème traité. Pour un problème de classification, nous voulons prédire une distribution de probabilité sur un ensemble de classes. Dans les problèmes de régression, cependant, nous voulons prédire une valeur spécifique plutôt qu'une distribution. Nous vous présentons les fonctions de perte de base, les plus couramment utilisées et proposons le type de problème auquel pourrai être associé pour donner un résultat optimal [51].

- Perte d'erreur quadratique moyenne (Mean Squared Error Loss) MSE: est calculée comme la moyenne des différences quadratiques entre les valeurs prédites et réelles. Le résultat est toujours positif, quel que soit le signe des valeurs prédites et réelles. Et une valeur parfaite est 0,0. La valeur de perte est minimisée, bien qu'elle puisse être utilisée dans un processus d'optimisation de maximisation en rendant le score négatif.
Il est conseillé de l'utiliser pour des problèmes de Régression.
- Perte d'entropie croisée (Cross-Entropy Loss) : référence à une perte logistique, chaque probabilité prédite est comparée à la valeur de sortie de classe réelle (0 ou 1) et un score est calculé qui pénalise la probabilité en fonction de la distance de la valeur attendue. La pénalité est logarithmique, offrant un petit score pour les petites différences (0,1 ou 0,2) et un score énorme pour une grande différence (0,9 ou 1,0). La perte d'entropie croisée est minimisée, où des valeurs plus petites représentent un meilleur modèle que des valeurs plus grandes. Un modèle qui prédit des probabilités parfaites a une entropie croisée est de 0,0.
Elle est optimale pour des problèmes de classification binaire, ou multiclassées.

6.7. Perceptron multicouche (Multilayer Perceptron) MLP

Le perceptron multicouche (MLP) relie plusieurs perceptrons (communément appelés neurones) ensemble dans un réseau, les neurones qui prennent la même entrée sont regroupés en une couche de perceptrons. Un MLP doit contenir une couche d'entrée et une de sortie et au moins une couche cachée. De plus, les couches sont également entièrement connectées (FC) ce qui signifie que la sortie de chaque couche est connectée à chaque neurone de couche suivante. En d'autres termes, un paramètre de poids est appris pour chaque combinaison de neurone d'entrée et neurone de sortie entre les couches [51].

6.8. Réseaux de Neurones convolutifs (Convolutional neural networks) CNN

Les réseaux convolutifs sont une forme particulière de réseau neuronal multicouche dont l'architecture des connexions est inspirée de celle du cortex visuel des mammifères. Leur conception suit la découverte de mécanismes visuels dans les organismes vivants [52]. Ces réseaux de neurones artificiels sont capables de catégoriser les informations des plus simples aux plus complexes. Ils consistent en un empilage multicouche de neurones, des fonctions mathématiques à plusieurs paramètres ajustables, qui prétraitent de petites quantités d'informations [53-56].

CNN est type spécifique de réseaux de neurones qui sont généralement composés des couches suivantes :

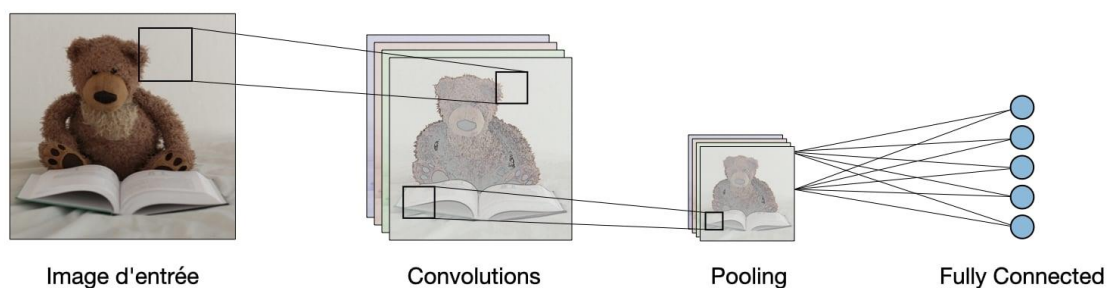


Figure 6 Différentes Couches du CNN

- **Couche convolutionnel (CONV)** : La couche convolutionnel (en anglais convolution layer) utilise des filtres qui scannent l'entrée suivant ses dimensions en effectuant des opérations de convolution. La sortie de cette opération est appelée feature map ou aussi activation map.
- **Pooling (POOL)** : La couche de pooling (en anglais pooling layer) est une opération de sous-échantillonnage typiquement appliquée après une couche convolutionnel. En particulier, les types de pooling les plus populaires sont le max et l'average pooling, où les valeurs maximales et moyennes sont prises, respectivement.

TYPE	Max Pooling	Average Pooling
But	chaque opération de pooling sélectionne la valeur maximale de la surface	chaque opération de pooling sélectionne la valeur moyenne de la surface
Spécificité	• Garde les caractéristiques détectées • Plus communément utilisé	• Sous-échantillonne la feature map • Utilisé dans LeNet

Tableau 2 Comparaison entre le max pooling et l'average pooling

- **Fully Connected (FC)** : la couche de fully connected (en anglais fully connected layer) (FC) s'applique sur une entrée préalablement aplatie où chaque entrée est connectée à tous les neurones. Les couches de fully connected sont typiquement présentes à la fin des architectures de CNN et peuvent être utilisées pour optimiser des objectifs tels que les scores de classe.

Le réseau de neurones convolutif sera détaillé davantage dans le chapitre 4.

6.9. Réseaux neuronaux à action directe (Feed-Forward Neural Networks)

Dans un réseau de neurones Feed-Forward, l'information ne se déplace que dans une direction, de la couche d'entrée, à travers les couches cachées, à la couche de sortie. L'information se déplace directement à travers le réseau. Pour cette raison, l'information ne touche jamais un nœud deux fois.

Feed-Forward Neural Networks n'ont aucune mémoire de l'entrée qu'ils ont reçue précédemment et son donc mauvais dans la prédiction de ce qui va suivre. Parce qu'un réseau feedforward ne considère que l'entrée actuelle, il n'a aucune notion d'ordre dans le temps. Ils ne peuvent tout simplement pas se souvenir de ce qui s'est passé dans le passé, sauf de leur entraînement [51].

6.10. Réseaux de neurones récurrents (Recurrent Neural Networks) RNN

Les réseaux récurrents (ou RNN pour Recurrent Neural Networks) sont des réseaux de neurones dans lesquels l'information peut se propager dans les deux sens, y compris des couches profondes aux premières couches. En cela, ils sont plus proches du vrai fonctionnement du système nerveux. Ces réseaux possèdent des connexions récurrentes au sens où elles conservent des informations en mémoire : ils peuvent prendre en compte à un instant t un certain nombre d'états passés. Pour cette raison, les RNNs sont particulièrement adaptés aux applications faisant intervenir le contexte, et plus particulièrement au traitement des séquences temporelles comme l'apprentissage et la génération de signaux. Les modèles

Apprentissage en profondeur

RNN sont surtout utilisés dans les domaines du traitement automatique du langage naturel et de la reconnaissance vocale [51].

Dans un RNN, l'information parcourt une boucle. Quand il prend une décision, il prend en compte l'entrée courante et aussi ce qu'il a appris des entrées qu'il a reçues précédemment.

Les avantages et inconvénients des architectures de RNN traditionnelles sont résumés dans le tableau ci-dessous :

Avantages	Inconvénients
• Possibilité de prendre en compte des entrées de toute taille. • La taille du modèle n'augmente pas avec la taille de l'entrée. • Les calculs prennent en compte les informations antérieures.	• Le temps de calcul est long. • Difficulté d'accéder à des informations d'un passé lointain. • Impossibilité de prendre en compte des informations futures un état donné.

Tableau 3 Avantages et inconvénients de l'architectures de RNN

Exemple : On a un réseau de neurones feed-forward normal et qu'on souhaite lui donner le mot "neurone" comme une entrée et qu'il traite le mot caractère par caractère. Au moment où il atteint le caractère "r", il a déjà oublié "n", "e" et "u", ce qui rend presque impossible à ce type de réseau de neurones de prédire quelle lettre viendrait ensuite.

Un réseau neuronal récurrent est capable de se souvenir exactement de cela, à cause de sa mémoire interne. Il produit une sortie, copie cette sortie et la boucle dans le réseau. Les réseaux neuronaux récurrents ajoutent le passé immédiat au présent.

6.11. Mémoire à long/court terme (Long Short Term Memory) LSTM

Les réseaux de mémoire à long/court terme sont une extension des réseaux neuronaux récurrents, qui étend leur mémoire. Par conséquent, il est bien adapté pour apprendre des expériences importantes qui ont des retards très longs entre les deux.

Les LSTM permettent aux RNN de se souvenir de leurs intrants sur une longue période de temps. C'est parce que les LSTM contiennent leurs informations dans une mémoire, ce qui ressemble beaucoup à la mémoire d'un ordinateur parce que le LSTM peut lire, écrire et supprimer des informations de sa mémoire.

Cette mémoire peut être vue comme une cellule gated, où gated signifie que la cellule décide de stocker ou de supprimer des informations, en fonction de l'importance qu'elle attribue à l'information. L'attribution de l'importance se fait à travers des poids, qui sont également appris par l'algorithme. Cela signifie simplement qu'il apprend avec le temps quelle information est importante et laquelle ne l'est pas.

6.12. Mémoire à long terme à court terme bidirectionnelle (Bidirectional Long Short Term Memory) BiLSTM (detaile de larchi)

Il ressemble exactement à son homologue unidirectionnel. La différence est que le réseau n'est pas seulement connecté au passé, mais aussi à l'avenir. Par exemple, les LSTM unidirectionnels peuvent être entraînés à prédire le mot «poisson» en alimentant les lettres une par une, où les connexions récurrentes dans le temps se souviennent de la dernière valeur. Un BiLSTM recevrait également la lettre suivante dans la séquence de la passe arrière, lui donnant accès à de futures informations. C'est pourquoi il est très utilisé dans le TAL.

7. Conclusion

En conclusion, le deep learning a été et est utilisé dans de nombreux domaines, y compris le TAL, il promet une révolution et une scalabilité a notre domaine qui est l'estimation du niveau de dépression . Nous explorerons dans le chapitre suivant certains des travaux effectués liés à l'estimation de niveau de dépression dans les réseaux sociaux avec les techniques du deep learning, ainsi que quelques études approfondies intéressantes qui nous ont aidés dans notre recherche.

Chapitre 03

La mesure du degré de dépression : état de l'art

Chapitre 03 : La mesure du degré de dépression : état de l'art

1. Introduction

Comme dans tout autre projet de recherche, l'étude des travaux antérieurs réalisés dans le même domaine et pour résoudre la même problématique joue un rôle très important dans l'utilisation de la bonne approche et la recherche de la solution la plus appropriée.

Par conséquent, dans ce chapitre, nous partagerons les points forts de notre revue de littérature malgré la difficulté rencontrée, car peu de travaux existent dans ce domaine qui est la mesure du degré de dépression.

Dans ce chapitre nous présentons les différents travaux effectués. Nous commençons tout d'abord par les troubles mentaux de façon générale pour par la suite entamer les recherches effectuées sur les troubles dépressifs majeurs afin de pouvoir aborder comment les réseaux sociaux ont été explorés pour détecter les personnes dépressives ceci mené aux travaux de la prédiction du niveau de dépression. À la fin, nous discuterons les différents résultats obtenus lors des travaux ancestraux sur l'estimation du niveau de dépression pour arriver à une conclusion afin de décrire dans le prochain chapitre notre modèle proposé.

2. Troubles mentaux

Selon [57] la maladie mentale, semble être un concept global et compréhensible, surtout quand il est basé sur la notion que le comportement de quelqu'un diffère radicalement des normes sociales acceptées. Les définitions de la maladie mentale sont également influencées par les idées, le climat politique et l'évolution récente de la société.

D'après [58, 59] une maladie mentale spécifique est considérée comme une maladie qui provoque un ensemble défini de symptômes ou de déficits fonctionnels. Ce point de vue repose sur un paradigme normatif qui suppose qu'il y a quelque chose comme "mentally healthy" des personnes qui n'ont pas de maladie mentale. Elle a également un impact sur le traitement et sur la façon dont les symptômes et les comportements problématiques sont compris par les cliniciens et les patients.

On utilise la description ou la définition communément acceptée comme indiqué par le American Psychiatric Association [60] comme suit : « Le trouble mental est un syndrome caractérisé par une perturbation cliniquement significative de la cognition, des émotions régulation ou comportement qui reflète un dysfonctionnement psychologique, biologique, ou les processus de développement qui sous-tendent le fonctionnement mental. Les troubles

mentaux sont généralement associés à une détresse importante dans la vie sociale, professionnelle ou autres activités ... »

Les troubles mentaux se présentent sous de nombreuses formes différentes, allant des symptômes d'anxiété temporaires aux expressions psychotiques. L'un des diagnostics les plus courants et les plus connus chez les adultes comme chez les plus jeunes, la dépression. La dépression dans sa forme clinique (pas nécessairement seulement les symptômes subcliniques) est considérée comme un trouble mental.

Nous avons pu conclure que les définitions de la maladie mentale varient considérablement d'un contexte à l'autre. Ce que l'on appelle la maladie mentale peut en fait faire partie d'un continuum normal et d'un autre type de comportement socialement inacceptable.

3. Troubles dépressifs majeurs (mettre année de lartcile)

Le trouble dépressif majeur MDD³ est une maladie mentale courante et fréquemment récurrente qui peut survenir à tout moment de la vie [60]. Elle est associée à une augmentation de la morbidité et de la mortalité dans le monde et a été classée comme la troisième cause principale parmi toutes les causes d'incapacité dans l'étude Global Burden of Disease Study 2015. Dans une méta-analyse de 18 études de prévalence et 5 études d'incidence, la prévalence annuelle et à vie et l'incidence annuelle du MDD étaient respectivement de 4,1%, 6,7% et 2,9%⁴.

Moins de 10% des personnes affectées dans le monde bénéficie de traitements connus et efficaces pour combattre leurs problèmes de troubles mentaux, du aux manques de ressources, de pénurie des soignants qualifiés, et c'est aussi du a la stigmatisation sociale liée aux troubles mentaux.

La charge de la dépression et des autres pathologies mentales est en augmentation dans le monde. Dans tous les pays, quel que soit leur niveau de revenu, les personnes déprimées ne sont pas souvent diagnostiquées correctement alors que d'autres, non affectées par ce trouble, se voient parfois prescrire des antidépresseurs par suite d'une erreur de diagnostic. Une résolution adoptée en mai 2013 par l'Assemblée mondiale de la Santé préconisait une réponse globale coordonnée au niveau des pays pour faire face aux troubles mentaux [61].

Ces dernières années, les communautés de santé mentale en ligne sont devenues des ressources de premier plan pour le soutien en santé mentale [62]. Ce soutien peut aller du soutien affectif (SE) au soutien informationnel (SI), en prenant souvent la forme d'empathie,

³ Major depressive disorder

⁴ Global Burden of Disease Study 2015 (<http://ghdx.healthdata.org/gbd-results-tool>)

de reconnaissance, de conseils ou d'une évaluation situationnelle sur divers enjeux comme la maladie mentale, la crise, la toxicomanie et la violence [63-65]. En fait, le soutien provenant de ces communautés a permis la réduction de la probabilité de pensées suicidaires [66]. En outre, en raison de la grande qualité du soutien fourni par ces OMHCS (Online Mental Health Communities). Les gens obtiennent souvent une conformité énorme avec une communauté par des moyens linguistiques : « Lorsque les gens interagissent, ils adaptent leur discours, leurs schémas vocaux et leurs gestes, pour s'adapter aux autres ». Plus précisément, les conventions d'une communauté sont généralement adoptées au fil du temps et se manifestent par des règles de communication ou des styles linguistiques établis par ses membres [67-70]. De tels accommodements sont liés à l'amélioration de la rétroaction sociale, à une solidarité accrue, à de meilleurs échanges sociaux et à des sentiments réciproques d'intimité [67]. Les OMHCS sont considérés comme un « refuge » ils permettent aux personnes d'exprimer des émotions désinhibées, de s'auto-dévoiler d'expériences stigmatisées et d'établir des relations de confiance avec leurs pairs [64-66, 71].

4. Exploration des réseaux sociaux pour détecter les personnes dépressives

Les médias sociaux sont devenus un élément indispensable de la vie de presque tout le monde. Selon [72], les médias sociaux étant n'importe quel site Web qui permet l'interaction sociale, ils ont considérablement amélioré la qualité de vie, apportant une grande efficacité.

Choudhury et tous [73]. Ont épluché la possibilité d'utiliser les réseaux sociaux pour identifier et analyser les problèmes dépressifs importants chez les humains. Grâce à leurs messages sur les réseaux sociaux, ils ont quantifié les crédits comportementaux en identifiant engagement social, sentiment, dialecte et styles sémantiques pour identifié leur mental désaccord.

O'Dea et tous [74] aussi affirme que Twitter est indispensable pour étudier l'état de bien-être psychologique , y compris la dépression dans la population, et a montré qu'il reconnaissait le niveau de dépression anxiété parmi les tweets liés au suicide.

Par ailleurs, Zhang et tous [75]. ont montré que si des individus présentant un risque élevé de suicide peuvent être reconnus par le biais d'un réseau en ligne comme un microblog, il est concevable d'actualiser un système d'intervention dynamique pour leur sauver la vie.

De nos jours, un certain nombre de chercheurs utilisent des groupes en ligne pour examiner les problèmes de bien-être psychologique. Nguyen et tous [76]. Ont utilisé l'apprentissage automatique et des stratégies statistiques pour séparer les messages en ligne entre les groupes dépressifs et témoins en utilisant le tempérament, les procédures psycholinguistiques et les sujets de substance retirés des postes créés par des individus de ces groupes.

Les recherches actuelles ont prouvé que les personnes souffrant d'une ou de plusieurs maladies mentales subiront probablement un effet boule de neige vers une autre, ce qui augmentera de façon exponentielle la probabilité de suicide d'après Chance et tous [77]. Selon Levine et tous [78]. Il y a un plus grand nombre de personnes qui se suicident par rapport à d'autres qui commettent un homicide, plus précisément, pour deux homicides, nous en avons trois personnes qui se suicident. Cela dit, Hall et tous [79], explique que la dépression a été la cause de plus des deux tiers des 30 000 suicides signalés au cours de la dernière année. De plus, d'autres études ont montré que la dépression non traitée est la première cause de suicide chez les jeunes, sachant que ce dernier lui-même est la troisième cause de décès chez les enfants selon Mann et tous [80]. Pour conclure, ces statistiques mettent en évidence certaines des réalités et des dommages potentiels causés par la dépression, nous amenant à convenir de l'urgence de prêter attention et de résoudre ce problème hautement prioritaire qui menace notre société.

L'étude [81] portait sur l'utilisation l'isolement social et des médias sociaux chez les jeunes adultes américains. L'étude disposée d'un échantillon de 1 787 jeunes de 19 à 32 ans représentatif à l'échelle nationale. Ils ont étudié l'utilisation par les participants des 11 plateformes de réseaux sociaux : Facebook, Twitter, Google+, YouTube, LinkedIn, Instagram, Pinterest, Tumblr, Vine, Snapchat et Reddi. L'étude a dénoncé que ceux qui ont visité n'importe quelle plateforme au moins 58 fois par semaine étaient trois fois plus susceptibles de se sentir socialement isolés que ceux qui ont utilisé les médias sociaux moins de 9 fois par semaine. En plus des sentiments d'isolement social et de dépression, les médias sociaux ont également été associés à l'image de soi. Une étude a révélé qu'une plus grande utilisation d'Instagram était associée à une plus grande auto-objection et préoccupation au sujet de l'image corporelle.

La technique [82] utilise l'apprentissage automatique pour classer automatiquement les patients anxieux, sentiments intérieurs et émotions en utilisant des classificateurs en fonction des modèles syntaxiques n-grammes. L'ensemble de données utilisé par des psychologues experts spécialement destiné à l'analyse des sentiments pour la santé mentale. Dans cet écrit, la classe traditionnelle de polarité « neutre » a été divisée en deux un sentiment double polarité (à la fois positive et négatif) et un sentiment ni positif ni négatif. Une polarité à quatre classes et binaire. Les expériences de prédiction et les résultats ont montré que ces dernières ont donné de meilleurs résultats, précision de 90% alors que la première a échoué avec seulement 78%.

Dans l'étude [83], différents modèles d'apprentissage profond ont été expérimentés dans cette étude (CNN, RNN et GRU), l'ensemble de données utilisé a été généré par la collecte des tweets de différentes pages. D'autres critères ont été pris en considération pour comparer les performances, ils ont examiné le modèle basé sur le caractère contre le modèle basé sur le mot et incorporation contre incorporations préentraînées. Les résultats ont montré que sur la base de mots GRU ont surclassé tous les autres modèles avec une précision de 98%. Un CNN

basée sur les mots a été considérée comme l'un des modèles les plus performants, car il a donné lieu à une précision de 97%.

Dans [84] ils abordent une approche de détection précoce du risque de dépression et l'anxiété sur les réseaux sociaux dans CLEF eRisk 2018. Pour les deux maladies mentales, les modèles combinent l'information TF-IDF (Term frequency-inverse document frequency). TF-IDF est calculé pour chaque mot, et les mots avec un faible score TF-IDF seront supprimés dans la phrase. Enfin, le classificateur de phrase est formé avec les phrases raffinées. Le reste de la séquence sera suivi en utilisant l'encodeur de phrases basé sur les réseaux de neurones convulsifs (CNN) pour identifier les articles écrits par des patients potentiels. L'évaluation officielle montre que le modèle atteint un F-score de 0,67 dans la détection de l'anxiété.

5. Mesure du degré de dépression sur les réseaux sociaux

Pour 2019, le CLEF e-Risk introduit pour la première fois un challenge visant à estimer le niveau de dépression à partir du fil des soumissions des utilisateurs [85, 86]. Pour chaque participant un ensemble de 20 fichiers, où chaque fichier contenant l'historique complet des publications d'un seul utilisateur. La tâche des participants consiste à remplir un questionnaire standard sur la dépression basé sur les preuves trouvées dans l'historique des publications. Les questionnaires sont dérivés de l'inventaire de dépression de Beck (BDI) [87], qui évalue la présence de sentiments comme la tristesse, le pessimisme, la perte d'énergie, etc.

Pour la détection de la dépression. Le questionnaire contient 21 questions (voir annexe). La tâche vise à explorer la viabilité d'une estimation automatique de la gravité des multiples symptômes associés à la dépression. Compte tenu de l'historique des écrits des utilisateurs qui varie de 30 jusqu'à 1511 postes par personne.

Les algorithmes proposés par les participants devaient estimer la réponse de chaque utilisateur à chaque question individuelle du BDI. La vérité fondamentale utilisée pour l'évaluation des réponses fournies par les participants à cette tâche était les questionnaires remplis par les utilisateurs des réseaux sociaux dont les écrits ont été fournis dans l'ensemble de données.

Les paramètres suivants Average Hit Rate (AHR), verage Closeness Rate(ACR), Average Difference between overall depression levels (ADODL), Depression Category Hit Rate (DCHR) ont été pris en compte pour l'évaluation des résultats [86], ils seront détaillés dans le chapitre 5.

Dans cette partie nous présentons quatre travaux sur la mesure de gravité des signes de dépression des participants au challenge CLEF e-risk 2019.

[88] ont développé un data-driven : qui est un modèle d'ensemble basé sur les données combinant des lexiques de sentiments, des informations mutuelles et des similitudes

d'intégration afin de surmonter le manque d'échantillons d'apprentissage. Ils proposent quatre modèles pour remplir automatiquement le questionnaire BDI en utilisant les messages Reddit de l'utilisateur qui sont : la polarité des mots, l'information mutuelle, la similitude sémantique et l'ensemble.

Dans le premier modèle qu'ils proposent, ils visent à tirer parti des polarités des mots pour classer d'abord les messages Reddit comme dépressifs en utilisant le lexique de subjectivité Multi Perspective Question Answering (MPQA) [89, 90] pour qu'après ils associent les postes les plus pertinents aux réponses du BDI.

En occurrence dans le second modèle d'informations mutuelles, ils tentent de créer un ensemble de données d'apprentissage à partir de Reddit pour classer les postes comme dépressifs ou non. Ils utilisent la mesure d'information mutuelle pour extraire les jetons pertinents des messages dépressifs [91]. Un classificateur de régression logistique a été formé pour classer les publications en dépressifs ou non dépressifs en utilisant les messages positifs et négatifs. Ensuite, les mots clés des catégories du BDI ont été développés à l'aide de WordNet [92] et utilisés pour baliser les messages classés positivement.

Dans le troisième modèle, ils ont proposé de trouver les messages d'utilisateurs les plus sémantiquement similaires aux réponses du questionnaire BDI afin d'estimer comment un utilisateur peut répondre au questionnaire. Étant donné le Word embeddings⁵, ils génèrent la représentation de chaque message d'utilisateur en calculant la moyenne des Word embeddings dans le message. Ils utilisent une approche similaire pour représenter les vecteurs de réponses au questionnaire, et font la moyenne des Word embeddings dans chaque réponse, puis calculent la similarité d'un message d'un poste d'utilisateur et d'une réponse au questionnaire en utilisant la similitude en cosinus⁶, ils utilisent le vecteur GloVe préentraîné Word embeddings [93] pour représenter les mots. Afin de filtrer les messages non pertinents, ils suppriment les messages qui ne sont pas similaires aux réponses du questionnaire BDI, par contre pour les autres, ils calculent la distance vectorielle de chaque message et réponse au questionnaire et choisissent la réponse la plus similaire, ils n'ont pas pris en considération l'ordre de mots et ne produisent probablement pas une bonne représentation pour les messages plus longs.

Dans le dernier modèle, ils ont testé deux approches d'ensemble : le micro et le macro-voting en utilisant les modèles de 1 à 3. Dans le micro-voting, les modèles ont été combinés au niveau du poste. Si plus de deux modèles classaient un poste comme positif pour la dépression et si deux modèles (majoritaires) ou trois (stricts) classaient le poste comme positif pour une catégorie, la catégorie est considérée comme positive pour cet utilisateur. Dans le macro

⁵ Plongement de mots

⁶ Cosine similarity en anglais : permet de calculer la similarité entre deux vecteurs à n dimensions en déterminant le cosinus de l'angle entre eux

voting, les résultats ont été combinés en utilisant la prédiction de catégorie moyenne des trois modèles précédents. Le run officiel (BiTeM run 0) a été généré en utilisant l'ensemble strict de micro-voting.

Au cœur du travail [88], ils disent que répondre au questionnaire BDI sans données d'apprentissage s'est avéré être une tâche difficile, avec un taux moyen de réussite de moins de 42% pour le premier système. En effet, pour certaines métriques, leur système a été superforme par une ligne de base de réponse naïve toute positive dans une classification binaire.

L'approche [94] consiste à utiliser des ensembles de données similaires pour fournir des entraînements étiquetés pour l'apprentissage automatique supervisé. Pour tirer parti de l'apprentissage par transfert [94], ils ont principalement utilisé des fonctionnalités extraites d'un modèle de langage appris via un entraînement non supervisé GPT-1 [95]. Pour l'extraction des fonctionnalités, ils utilisent GPT-1. Le modèle GPT préentraîné a été *fine_tuned*⁷ sur le texte fourni par les 20 sujets de l'ensemble de données eRisk. Ils utilisent pour extraire les fonctionnalités l'outil LingWistic Inquiry Word Count (LIWC) [96], et ils ont utilisé toutes les catégories LIWC disponibles, ce qui donne 70 fonctionnalités par publication.

Pour prédire les réponses BDI, ils utilisent deux manières : la manière non supervisée et la supervisée.

Premièrement, de manière non supervisée, ils utilisent deux approches : dans la première approche, ils visent à utiliser les fonctionnalités du vecteur au niveau de la phrase (*sentence_level* en anglais) de chaque poste d'utilisateur qui a été fournie par GPT-1. Puis agréger les fonctionnalités du vecteur au niveau de la phrase en vecteurs de fonctionnalité au niveau de l'utilisateur, toujours en utilisant GPT-1, ils génèrent un vecteur de fonctionnalité pour les réponses aux questions BDI, pour remplir le questionnaire pour un utilisateur. Pour chaque question, ils ont calculé le Similitude de cosinus entre le vecteur caractéristique de l'utilisateur et le vecteur caractéristique de chacune des réponses possibles à cette question. La réponse qui a donné la valeur de similitude cosinus la plus élevée a été sélectionnée pour être la réponse de l'utilisateur à la question.

Dans la deuxième approche, ils n'ont pas agrégé les caractéristiques du vecteur au niveau de la phrase en un seul vecteur de caractéristiques au niveau de l'utilisateur. Au lieu de cela, pour prédire la réponse d'un utilisateur à chacune des questions, ils ont calculé la similitude cosinus de chaque réponse possible à toutes les phrases d'un utilisateur. Ensuite, ils ont choisi la réponse qui avait la plus grande similitude de cosinus à l'une des phrases que cet utilisateur n'avait jamais écrit. Ceci est l'approche du voisin le plus proche qui demande quelle réponse

⁷ *Fine_tuned* un modèle prédictif d'apprentissage automatique est une étape cruciale pour améliorer la précision des résultats

possible à la question est la plus proche de ce que le sujet a déjà écrit dans l'espace des caractéristiques du GPT-1.

Enfin, de manière supervisée, ils ont utilisé les données d'une étude que l'un des membres de l'équipe a effectuées au cours de sa thèse de doctorat, pour la phase d'entraînement, ils ont utilisé la machine à vecteurs de support SVM entraîné, avec des noyaux linéaires, et la régularisation L2 en utilisant l'ensemble de données étiqueté.

Ils ont utilisé le modèle résultant pour générer des réponses prédictives à une question pour chaque utilisateur, et répéter ce processus pour chaque question du BDI.

Ils ont également testé une deuxième approche supervisée qui utilisait les caractéristiques du GPT-1 et AutoSklearn. Sur la base de leur approche non supervisée, ils ont testé des caractéristiques qui représentent la moyenne d'un utilisateur les caractéristiques du GPT-1 et les relations les plus étroites entre chaque réponse possible et les écrits d'un utilisateur. Pour ces relations, ils ont utilisé la distance euclidienne minimale et la corrélation maximale.

Pour évaluer leurs résultats [94] trouvent que répondre au questionnaire BDI même avec des données d'apprentissage s'est avéré être une tâche difficile, avec un taux de réussite moyen de moins de 24%. En effet, pour certaines métriques, leur système était surperformé par leur manière supervisée.

Les auteurs de [97] ont élaboré une nouvelle mesure basée sur l'utilisation de la langue durant leur participation à CLEF ERISK 2019. Pour identifier la présence de dépression et évaluer son niveau de gravité, ils ont formé un classificateur de texte binaire SS3 [44] pour détecter les utilisateurs « déprimés », puis ils utilisent les valeurs de confiance pour estimer le niveau de dépression clinique de l'utilisateur en remplissant le questionnaire BDI.

Ils ont entraîné SS3 en utilisant l'ensemble de données pour la tâche de détection de dépression de eRisk 2018 [59] en raison de l'absence de données d'apprentissage. De plus, ils ont utilisé les mêmes hyperparamètres utilisés dans [58], c'est-à-dire $\lambda = \rho = 1$ et $\sigma = 0,455$, pour lesquels SS3 a montré des performances de pointe en matière de détection de dépression.

Donc, tout d'abord, SS3 a traité l'intégralité de l'historique d'écriture en calculant la valeur de confiance, puis la valeur finale (0,941) a été utilisée pour prédire la catégorie de dépression, « dépression sévère » ($c = 3$). Enfin, le niveau de dépression a été calculé par le mapping en région / catégorie c , une pour chaque catégorie de dépression BDI ($c \in \{0, 1, 2, 3\}$). À ce stade, ils ont transformé la sortie de SS3 à partir d'un vecteur bidimensionnel en BDI niveau de dépression, mais ils n'ont pas encore découvert comment répondre réellement aux 21 questions du questionnaire BDI en utilisant ce niveau de dépression, car ils décident que pour tous les utilisateurs dont le niveau de dépression est inférieur ou égal à 0, toutes les questions BDI ont été répondues par un 0. Pour les autres utilisateurs, ils ont appliqué des méthodes

différentes puisque chaque équipe participante était autorisée à utiliser jusqu'à cinq modèles différents (appelés «runs») pour effectuer la tâche. Ainsi, ils utilisent cinq méthodes différentes pour accomplir cette tâche, comme décrits ci-dessous :

UNSLA : en utilisant le niveau de dépression prédit, leur modèle a rempli le questionnaire réponse en répondant au numéro de réponse sur chaque question. Si cette division avait un reste, les points restants étaient dispersés au hasard de sorte que la somme de toutes les réponses correspondait toujours au niveau de dépression prédit donné par SS3.

UNSLB : ici, seule la catégorie prédite a été utilisée, leur modèle a rempli le questionnaire de manière aléatoire de telle sorte que le niveau de dépression final correspondait toujours à la catégorie prédite.

UNSLC : proche de l'*UNSLA*, la réponse à ce nombre uniquement sur les questions pour lesquelles un « indice textuel » pour une réponse possible a été trouvé dans les écrits de l'utilisateur, et a répondu de manière aléatoire et uniforme.

UNSLD : le même que le précédent, mais sans utiliser les « conseils textuels », répondant toujours à chaque question de manière aléatoire et uniforme.

UNLSE : le même que le précédent, mais cette fois sans utiliser de distribution uniforme.

[97] disent que répondre au questionnaire BDI sans données d'apprentissage s'est révélé être une tâche difficile. Ils ont obtenu les meilleurs AHR (41,43%) et ACR (71,27%), et les deuxièmes meilleurs ADODL (80,48%) et DCHR (40%), les meilleurs DCHR et ADODL ont été obtenus par CAMH, l'utilisation de «textuel », dans UNSLC, a vraiment amélioré le Taux de réussite moyen (AHR) mais sans impact sur les autres mesures, UNSLE était considérablement la meilleure méthode pour estimer le niveau de dépression global puisque son ADODL prend des valeurs dans une fourchette bien au-dessus des autres.

[98] vise à explorer la viabilité de l'estimation automatique de la gravité des symptômes multiples associés à la dépression [86]. L'idée proposée consiste à faire un mixe d'approche qui combine l'apprentissage automatique avec la psycholinguistique et le schéma comportemental.

Leur approche pour résoudre cette tâche était basée sur des règles et chaque règle a été modélisée avec des références à plusieurs schémas comportementaux et psycholinguistiques qui sont connus pour être associés à l'état de dépression.

Ils témoignent que la taille réduite de l'ensemble de données, en termes d'utilisateurs et le petit nombre d'écrits par utilisateur, ainsi que l'absence d'un ensemble d'apprentissages a conduit au choix d'une approche fondée sur des règles plutôt que l'utilisation d'algorithmes d'apprentissage automatiques.

De plus, ils ont exploré la corrélation entre certaines questions par diviser les 21 questions en 6 groupes (Depression, Guilt, Appetite, Anxiety, Fatigue, Sleep) toutes les questions appartenant à un groupe donné ont été notées avec la même réponse ou valeur numérique.

Pour chaque catégorie, un score a été calculé pour chaque utilisateur comme une valeur normalisée du nombre d'occurrences des caractéristiques considérées pour chaque catégorie par rapport au nombre total d'occurrences des mêmes caractéristiques de la base de données. Ces scores ont ensuite été normalisés à l'intervalle [0,3] sur la base d'un seuil extrait des histogrammes d'occurrences.

Ils ont pu récolter un score de 34,5% en AHR tant dis que le meilleur score obtenu était 41,43 , en ACR ils ont obtenu un taux de 66,43% face au meilleur taux de 72,27%, pour l'ADODL le score aussi était proche il y avait une différence de 3,33% , pour le DCHR ils ont obtenu 25% tant dis que le meilleur score était de 45%.

6. Discussion

Dans le tableau suivant nous allons comparer les différents travaux sur l'estimation du degré de dépression sur les réseaux sociaux : (critique methode qui ont donner meilleure resultat)

Etude	Catégorie de l'analyse	Type d'analyse	Classification	Niveau d'analyse
[88]	Prédiction	Détection des émotions	Approche basé sur des règles	Niveau des aspects
[94]	Prédiction	Détection des émotions	hybride	Niveau des aspects
[97]	Prédiction	Analyse de sentiment a base d'aspect	Approche basé sur des règles	Niveau des aspects
[98]	Prédiction	Analyse de sentiment a base d'aspect	Approche automatique	Niveau des aspects

Tableau 4 Comparaison entre les travaux de mesure du niveau de dépression

La première chose que nous avons remarquée lors de notre revue de la littérature est que le sujet en lui-même est récent et on ne trouve pas beaucoup de travaux sur la mesure du degré de dépression sur les réseaux sociaux, et que la majorité de ces travaux sont des personnes qui ont participé au challenge E-risk de CLEF 2019. Chacun des participants a choisi un modèle différent. Dû au manque de l'ensemble des données la tâche s'est avérée difficile, certain participant témoigne même avoir eu des résultats non crédibles proches du hasard.

Nous avons déduit depuis l'évaluation de performances de eRisk 2019 que l'approche proposée basée sur des règles avec une analyse de données niveaux aspect sont celle qui ont donné les meilleurs scores.

7. Conclusion

Dans ce chapitre, nous avons présenté les travaux qui ont été menés avec la même problématique que la nôtre. Nous avons abordé les études les plus pertinentes et les plus influentes en termes de construction de notre modèle proposé et des bons critères.

La prochaine partie de ce mémoire sera consacrée à partager tous les détails sur notre modèle proposé et le processus que nous avons suivi pour le concevoir, le construire et le tester.

Chapitre 04

Conception d'un modèle de mesure du degré de dépression à partir des réseaux sociaux

Chapitre 04 : Conception d'un modèle de mesure du degré de dépression à partir des réseaux sociaux.

1. Introduction

Après avoir fait le tour sur les travaux connexes réalisés dans ce domaine, nous présentons dans ce chapitre l'architecture de notre modèle pour mesurer le niveau de dépression à partir des réseaux sociaux. Notre modèle est constitué de quatre parties. Premièrement, nous traitons les données stockées dans des fichiers « Subject », deuxièmement nous traitons les données d'apprentissage « Training dataset » et les données de tests « Test dataset », par la suite nous passons par l'apprentissage de trois modèles CNN , BiLSTM et CNN_BiLSTM, et en dernier nous évaluons nos modèles par le jeu de données « Test dataset ».

2. Architecture de notre modèle

Les quatre étapes essentielles de notre modèle sont les suivantes :

- Étape 1 : Extraction de données (Training dataset, Test dataset).
- Étape 2 : Prétraitement et vectorisation.
- Étape 3 : Apprentissage des modèles.
- Étape 4 : Test.

Dans les autres travaux, les chercheurs ont fait des études similaires pour résoudre ce problème en utilisant les méthodes SVM, SS3 ou GPT-1, mais selon nos recherches, aucune équipe scientifique n'a travaillé avec les méthodes d'apprentissages en profondeur telles que CNN ou bien BiLSTM dans le même but que le nôtre. L'utilisation de ces méthodes engendre la plus grande partie de notre modèle.

Dans la figure suivante (**figure 6**), nous avons schématisé les principales étapes constituant notre modèle de mesure du niveau de la dépression.

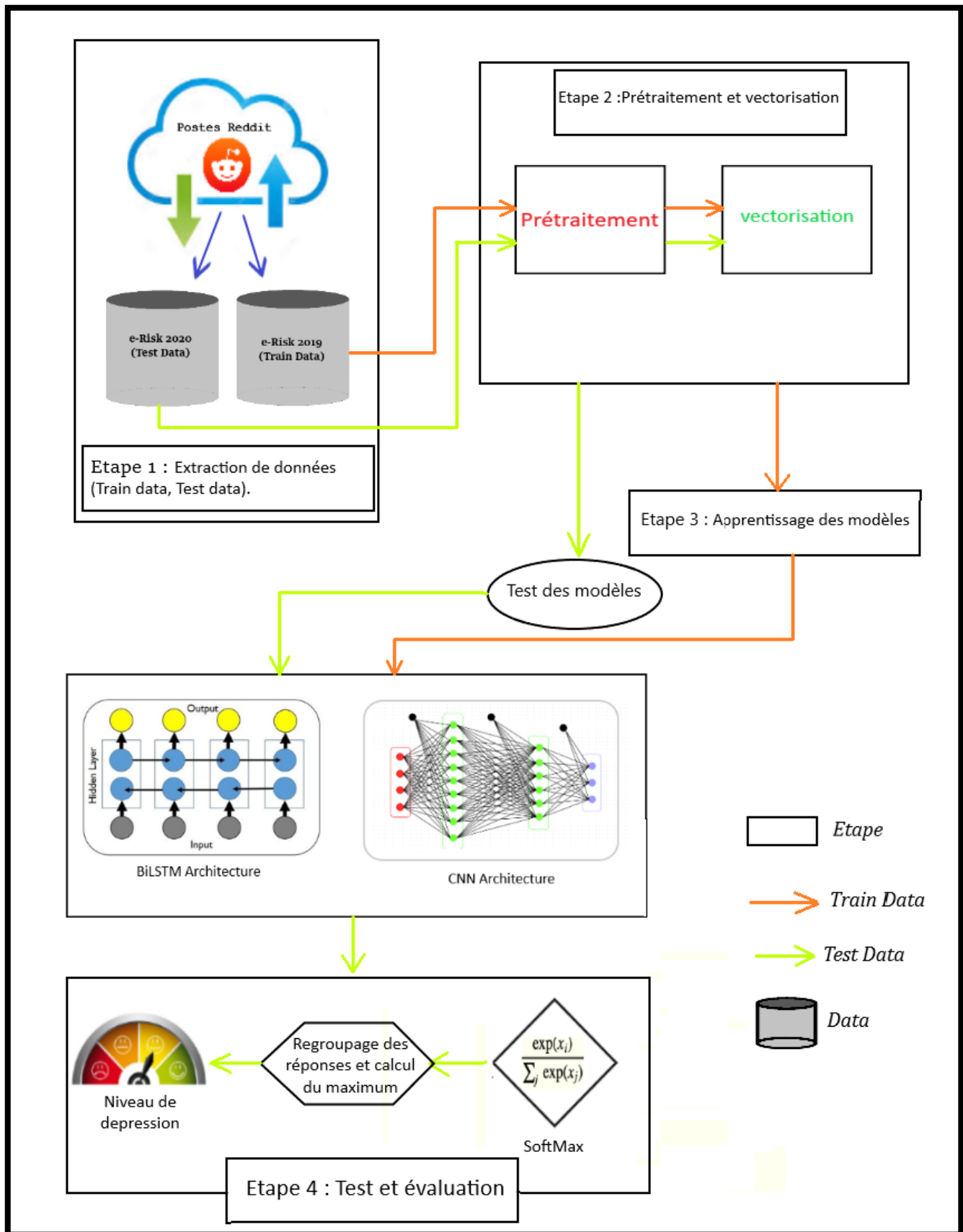


Figure 7 Architecture de notre modèle de mesure du niveau de la dépression

3. Extraction de données (Training data, Test data)

Nous utilisons les données d'apprentissage et les données de test proposées par eRisk 2020. Le but est de prédire les réponses du questionnaire de « BDI ». Plus de détails sur ces données sont présentées dans le chapitre 5.

eRisk 2020 met à la disposition des participants les données de test constituées de 70 fichiers XML qui correspondent à l'historique des publications des utilisateurs sur Reddit. Le dataset de eRisk 2019 est aussi fourni pour l'apprentissage avec un total de 20 fichiers et un fichier TXT (Depression Questionnaires_anon.txt) contenant les 21 réponses réelles du questionnaire de « BDI ». Comme le montre la figure (**Figure 7**), ces données sont stockées dans des fichiers « XML ».

```
<INDIVIDUAL>
<ID>subject436</ID>
<WRITING>
<TITLE> Hospital or no? </TITLE>
<DATE> 2018-07-09 05:18:12 </DATE>
<INFO> reddit post </INFO>
<TEXT> [removed] </TEXT>
</WRITING>
<WRITING>
<TITLE> D* on a 12 hour road trip...suggestions? </TITLE>
<DATE> 2018-06-10 14:22:03 </DATE>
<INFO> reddit post </INFO>
<TEXT> I'm traveling halfway across the country by car today with
traveling with woke up v*ing this morning, so I rode in the car h
and now that I'm in the car I'm feeling like I need to again, and
or worse, more of the adults start v*ing as well. I need suggesti
```

Figure 8 Structure des fichiers XML du Dataset de eRisk 2020

L'objectif principal de cette étape est de faire la combinaison entre les fichiers XML et le fichier TXT d'eRisk 2019 pour les regrouper dans un seul fichier CSV (**figure 7**) qui sera utilisé pour faire l'apprentissage de nos modèles, et extraire toutes les publications de chaque utilisateurs d'eRisk 2020 et les mettre dans des fichiers CSV individuellement (**figure 8**).

Personne	TEXT	Sadness	Pessimism	Past Failure	Loss of Pleasure	Guilty Feelings	Punishm
subject2341	if you need to talk to someone its not a bad thing everyone does and im sure at least some	1	2	3	2	3	
subject2341	its not really about fate or anything like that what i mean is your life is exactly the same as k	1	2	3	2	3	
subject2341	i beleive in determinism but i think of it a little differently than just you have no control bas	1	2	3	2	3	
subject2341	you just need to explain to her that too much advice and help stresses you out and that you	1	2	3	2	3	
subject2341	a covalent bond is when 2 or more orbitals overlap but electrons tend to stay in the orbitals	1	2	3	2	3	
subject2341	fallout new vegas	1	2	3	2	3	
subject2341	in organic chemistry reduction is when you gain something that increases charge density wi	1	2	3	2	3	
subject2341	it shouldnt ruin your friendship with either of them this could be a really good thong hones	1	2	3	2	3	
subject2341	im sorry im not sure what you need help with is this necessarily a bad thing	1	2	3	2	3	
subject2341	dont do that i know its hard to keep going but cutting your life short isnt the answer if your	1	2	3	2	3	
subject2341	ive had a lot of friends who talked to therapists most of them said that it helped them its gc	1	2	3	2	3	
subject2341	btw what do you mean making yourself useful	1	2	3	2	3	

Figure 9 Structure d'un fichier d'apprentissage csv

TEXT								
answers dont match the questions college education has no gender as far as i know but my incom								
how is once a year the lowest possible amount to visit a health professional in a year								
mark williams was ranked 16th in 1997 and i think he had never qualified before that								
i would specialize in birds reptiles and insects as its way too cold here it would either be mainly								
the goth is here to stay								
i won a summer job at a tech competition sadly the job turned out to be marketing and as i am a								
women who drink and dress provocatively deserve what they get what am i supposed to answer								

Figure 10 Structure d'un fichier de test csv

4. Prétraitement et vectorisation

Cette étape passe par deux phases :

- Premièrement le **PRETRAITEMENT**.
- Deuxièmement la **VECTORISATION**.

4.1. Prétraitement

Cette phase a pour objectif de réduire une publication d'un utilisateur à une série de mots les plus représentatifs du texte, et la standardise afin de rendre son usage plus facile, elle extrait ainsi l'information contenue dans la publication.

Les techniques que nous avons utilisées dans cette phase sont les suivantes :

- Passage en minuscule : nous transformons les majuscules en minuscules par ce que les étapes suivantes sont très sensible à la casse.
- Suppression de URL : les URL ne contiennent pas d'information sur le sentiment des publications des utilisateurs alors elles sont supprimées.
- Suppression des nombres et des chiffres : les nombres et les chiffres ne servent à rien dans l'analyse des sentiments donc nous les supprimons.
- Suppression de la ponctuation : la ponctuation n'a aucun effet sur les sentiments des utilisateurs alors elle est enlevée.

- Retrait des stopwords : les mots vides ou stopwords sont nocifs pour l'analyse des sentiments, ces mots n'apportent aucune information, et sont en général très courant et présent dans la plupart des publications, alors nous les retirons suivant une liste de stopwords prédéfinis.

Le traitement automatique du langage naturel vise à créer des outils de traitement de la langue naturelle pour diverses applications.

Figure 11 Exemple d'une phrase avec Stopwords

traitement automatique langage naturel vise créer outils
traitement langue naturelle diverses applications

Figure 12 Phrase de la figure 10 sans stopwords

- La Lemmatisation : La lemmatisation a pour but de trouver la forme canonique du mot, ainsi elle permet de regrouper des mots d'une même famille qui partagent le même suffixe lexical. Elle doit se faire après la transformation des lettres majuscules en minuscules et avant la tokenisation car les mots présents avant et après sont importants pour déterminer la nature du mot.

traiter automatique langue naturel viser créer outil
traiter langue naturel divers application

Figure 13 Phrase de la figure 11 après lemmatisation

4.2. Vectorisation

Cette deuxième phase utilise ces mots pour représenter chaque publication sous forme de vecteur. Les vecteurs obtenus sont interprétables par la machine ce qui permet l'apprentissage de nos modèles.

Les principales étapes de cette seconde phase sont :

- Tokenisation : il s'agit de décomposer une publication en un ensemble de mots, nommés jetons. La figure (**Figure 13**) présente la tokenisation pour la phrase d'exemple : « le traitement automatique du langage naturel vise à créer des outils de traitement de la langue naturelle pour diverses applications ».

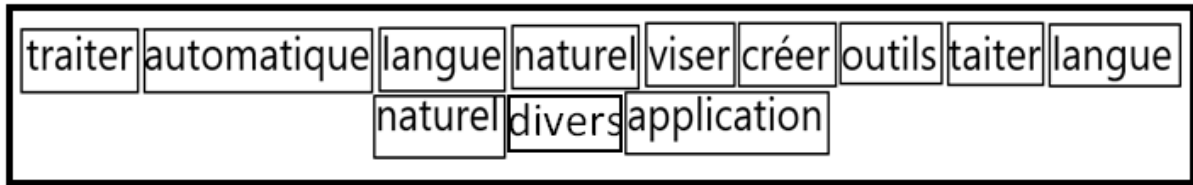


Figure 14 Tokens de la phrase de la figure 12

- Padding (Rembourrage) : cette étape est obligatoire avant de faire la représentation vectorielle, elle consiste à créer des séquences de mots de la même taille en fixant une taille maximale, les publications de taille inférieure à celle de la taille maximale seront remplis par des zéro, et les publications de taille supérieure seront tronquées.

Exemple :

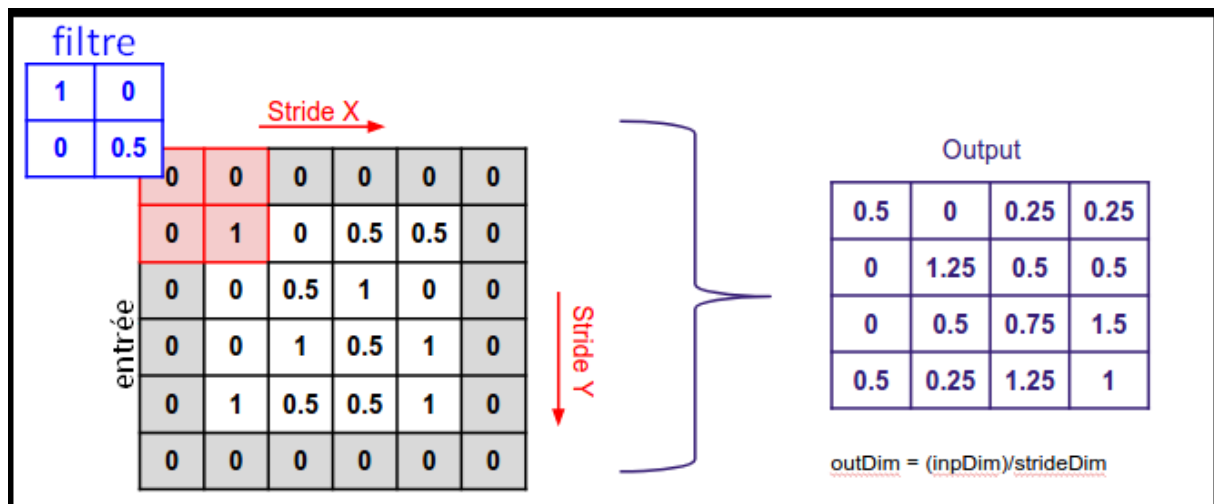


Figure 15 Exemple sur l'utilité du padding

- Word Embedding « plongement de mots » : c'est la représentation vectorielle d'un mot particulier d'un document par des nombres réels, projetés dans l'espace vectoriel (**Figure 15**) les uns à côté des autres. Cette technique nous permet de capturer le contexte d'un mot dans un document, la similarité sémantique et syntaxique, et même la relation avec d'autres mots.

Cette représentation vectorielle améliore les performances des algorithmes d'apprentissage dans le domaine du traitement automatique de la langue.

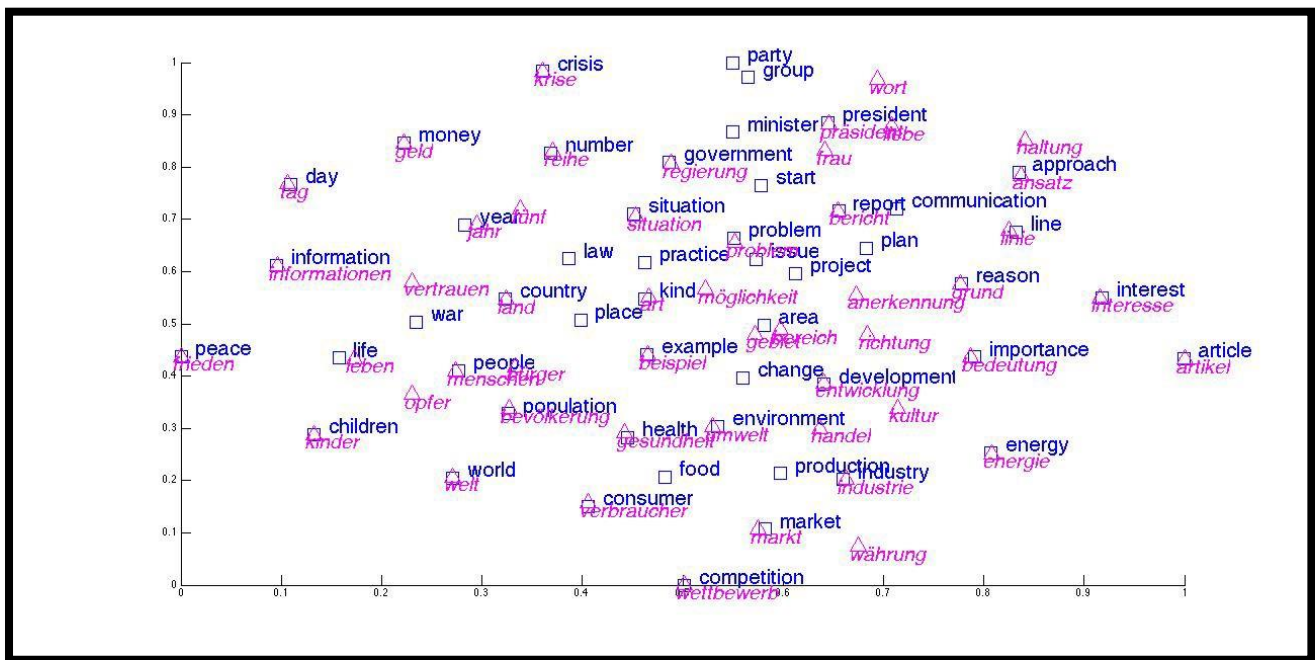


Figure 16 Exemple d'une représentation vectorielle

- Word2Vec : c'est un algorithme de word embedding, il repose sur un réseau de neurones à deux couches, ce dernier possède deux architecture neuronales appelées CBOW et Skip-Gram [99].
 - a) Modèle Skip-Gram : Le modèle skip-gram permet de prédire les mots environnants en utilisant le mot courant dans la séquence. Le score de classification des mots environnants est basé sur la relation syntaxique et les occurrences avec le mot central. (**Figure 13**), dans notre Word2Vec Skip-Gram y figure .
 - b) Modèle CBOW : Le modèle CBOW (**C**ontinuous **B**ag-**O**f-**W**ords) incite à prédire le mot courant à partir des mots environnants. (**Figure 13**)

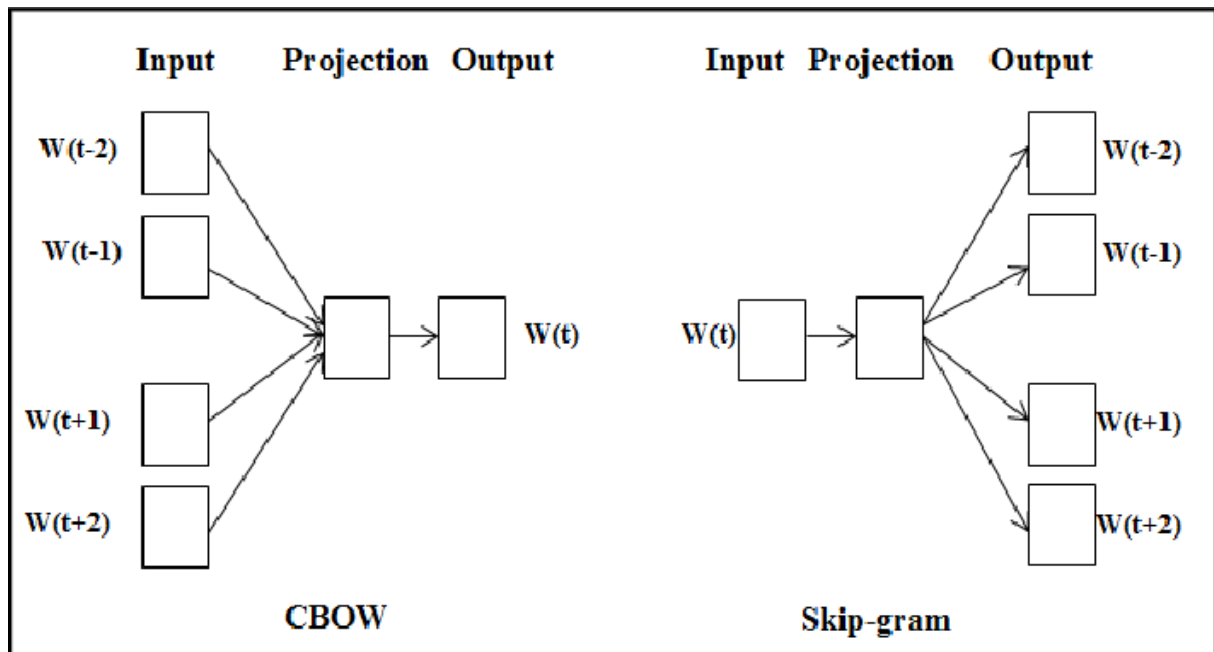


Figure 17 Modèle Skip-gram et CBOW [1]

5. Apprentissage des modèles

5.1. Modèle réseau de neurones à convolution « Convolutional Neural Networks CNN »

CNN est une classe de réseau de neurones à feed-forward profond et utilise un perceptron multicouche pour un pré-entraînement minimal. CNN initialement développé pour la classification d'images et les problèmes de vision par ordinateur. Récemment, il a été appliqué dans diverses tâches de TAL. Lorsque CNN est appliqué au texte au lieu à des images, il est nécessaire d'appliquer une représentation de tableau d'une dimension du texte.

Les résultats obtenus dans le TAL sont impressionnants en utilisant des vecteurs de mots pour construire la matrice d'entrée du modèle, le texte est traité de la même manière que les images, à la fois pour l'extraction de caractéristiques et la classification, et depuis le CNN est devenu l'un des réseaux de neurones les plus utilisés en TAL.

Au cœur du réseau neuronal convolutionnel se trouve plusieurs couches, nous citerons ci-dessous les couches que nous avons utilisées dans notre modèle :

5.1.1. Convolution

La couche convolutionnelle qui donne son nom au réseau. Quand on lui présente une nouvelle image, le CNN ne sait pas exactement si les caractéristiques seront présentes dans l'image ou où elles pourraient être. Le but de cette couche est donc de trouver dans toute l'image et dans n'importe quelle position des caractéristiques, pour ce faire il fait un filtrage à une entrée pour créer une carte de caractéristiques qui indique où la caractéristique a été trouvée dans l'image, l'opération réalisée est appelée une convolution, ce processus complet de convolution est

répété pour chaque caractéristique existante de l'image pour former le résultat qui est un ensemble d'images filtrées, dont chaque image correspond à un filtre particulier. En traitement automatique du langage, la convolution peut être considérée comme une fonction de fenêtre glissante (Sliding window) que nous appliquons à une matrice afin d'extraire des caractéristiques. Cette fenêtre est connue sous le nom de filtre ou de détecteur de caractéristiques (feature detector).

Les entrées de notre tâche sont des phrases représentées sous forme de matrice. Chaque ligne de la matrice correspond à un jeton, dans notre cas c'est un mot. Il peut être un caractère dans d'autres cas. Autrement dit, chaque ligne est un vecteur qui représente un mot. En règle générale, ces vecteurs sont des imbrications de mots (représentations de faible dimension) comme word2vec qu'on a utilisés au lieu de Glove, mais ils peuvent également être des vecteurs ponctuels qui indexent le mot dans un vocabulaire. Pour une phrase de 5 mots utilisant un embedding de 300 dimensions, nous aurions une matrice 5×300 comme entrée. C'est notre «image» en quelque sorte.

Dans la vision, nos filtres glissent sur les patches locaux d'une image, mais en TAL, nous utilisons généralement des filtres qui glissent sur des lignes complètes de la matrice (mots). Ainsi, la largeur «width» des filtres est généralement la même que la largeur de la matrice d'entrée. La hauteur «height», ou la taille de la région «region size», peuvent varier. En réunissant tout ce qui précède, un réseau neuronal convolutif pour le TAL peut ressembler à ceci :

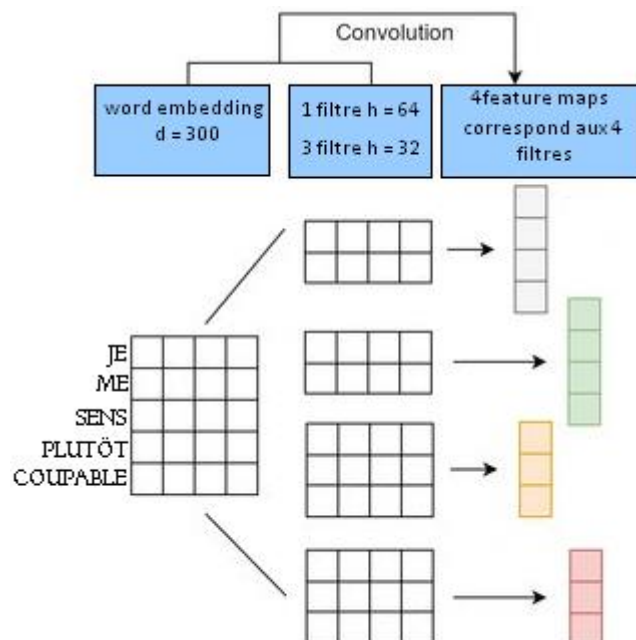


Figure 18 Un réseau neuronal convolutif pour le TAL

Pour construire une architecture CNN, quelques paramètres doivent être fixés avant de commencer. Pour appliquer le filtre au premier élément d'une matrice qui n'a pas d'éléments voisins en haut et à gauche, nous avons utilisé un remplissage nul appelé zéro padding. Tous les éléments qui sont en dehors de la matrice sont considérés comme nuls. En faisant cela, on

peut appliquer le filtre à chaque élément de la matrice d'entrée et obtenir une sortie plus grande ou de taille égale. L'ajout du zéro padding est également appelé convolution large « wide convolution », et ne pas utiliser de zéro padding serait une convolution étroite « narrow convolution ». Lorsqu'on a un grand filtre par rapport à la taille d'entrée il est conseillé d'utiliser wide convolution.

La formule générale pour calculer la sortie : $n_{out} = (n_{in} + 2 * n_{padding} - n_{filter}) + 1$

Dans notre modèle, nous avons appliqué quatre couches de convolution à une dimension. Le nombre de couches de convolution qu'on a choisi a été décidé de façon intuitive, car plus le nombre de couches est grand plus le modèle sera profond, aussi le nombre de couches qu'on choisit dépend des données, de leur taille. En outre pour choisir le nombre de couches c'est selon la problématique. Pour la dimension, nous avons choisi des convolutions à une dimension, car nous traitons du texte.

5.1.2. Les couches poolings (mise en commun)

Un aspect clé des réseaux de neurones convolutionnels est la mise en commun des couches (pooling), la propriété du pooling est qu'il fournit une matrice de sortie de taille fixe, qui est généralement requise pour la classification. Par exemple, si on a 1 000 filtres et qu'on veut appliquer une mise en pool maximale à chacun, on obtiendra une sortie de 1 000 dimensions, quelle que soit la taille de nos filtres ou la taille de notre entrée. Cela nous permet d'utiliser des phrases de taille variable et des filtres de taille variable, mais obtient toujours les mêmes dimensions de sortie pour alimenter un classificateur.

Le pooling réduit la quantité d'informations générées par la couche convolutionnel pour chaque entité et conserve les informations les plus essentielles (le processus des couches convolutionnels et de mise en commun se répète généralement plusieurs fois), il permet aussi de diminuer la charge de calculs. Par exemple : Si on considère chaque filtre comme la détection d'une caractéristique spécifique, comme la détection si la phrase contient une négation comme « pas heureux » . Si deux mots apparaissent quelque part dans la phrase, le résultat de l'application du filtre à cette région donnera une valeur élevée, mais une petite valeur dans d'autres régions. En effectuant l'opération max, on conserve des informations sur l'apparition ou non de la fonction dans la phrase, mais on perd des informations sur l'endroit exact où elle est apparue. Les informations sur la localité ne sont pas vraiment utiles, c'est un peu similaire à ce que fait un modèle de sac de n-grammes. On perd des informations globales sur la localité (où dans une phrase quelque chose se passe), mais on conserve des informations locales capturées par nos filtres, comme « pas heureux » étant très différent de « heureux pas ».

5.1.3. Unités Rectifié Linéaire (ReLU)

L'opération consiste à chaque fois qu'il y a une valeur d'entrée négative, elle la remplace par un 0. Ainsi, elle empêche les valeurs apprises de rester coincée autour de 0 ou d'exploser vers l'infinie.

5.1.4. La couche fully-connected (FC)

Elle est généralement utilisée comme la dernière couche d'un réseau de neurones convolutif, contrairement à la mise en commun (pooling) et la convolution, c'est une opération globale. Elle prend la contribution des étapes d'extraction des fonctionnalités et analyse globalement la sortie de toutes les couches précédentes. Par conséquent, Elle fait une combinaison non linéaire de caractéristiques sélectionnées, qui sont utilisées pour la classification des données [100].

Pour plus de précision, nous spécifions les sous-catégories de la couche FC :

- **La couche d'entrée entièrement connectée (Fully connected input layer)** : aplatis (Flatten) transformer les sorties générées par les couches précédentes en un seul vecteur qui peut être utilisé comme entrée pour la couche suivante.
- **Couche entièrement connectée (Fully Connected Layer)** : applique des pondérations à l'entrée générée par l'analyse d'entités pour prédire une étiquette précise.
- **La couche de sortie entièrement connectée (Fully connected output layer)** : génère les probabilités finales pour déterminer une classe pour l'image.

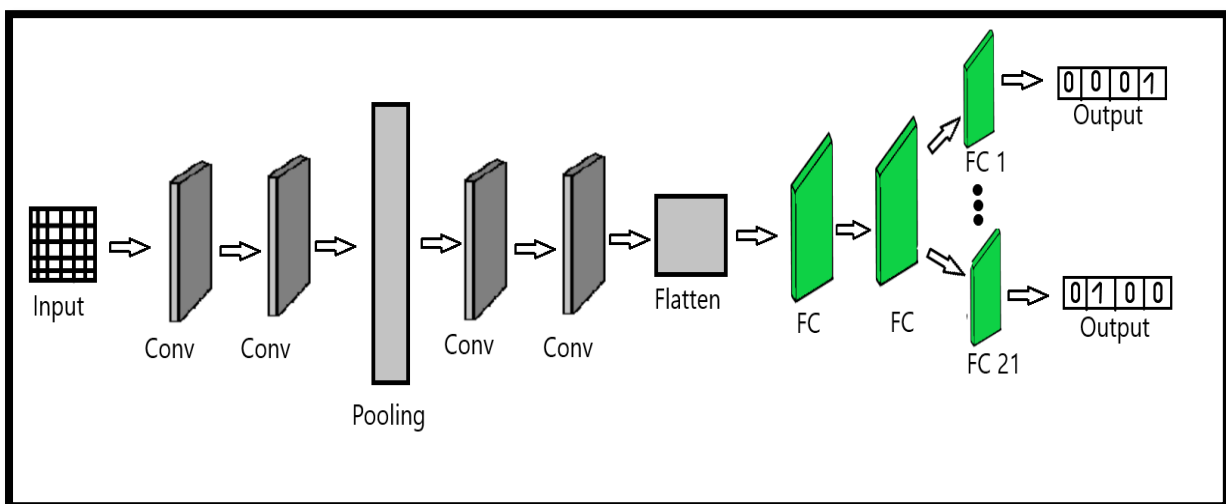


Figure 19 Architecture de notre modèle CNN

5.2. Modèle biLSTM (BI-DIRECTIONAL LONG SHORT TERM MEMORY)

Comme en TAL, parfois pour comprendre un mot, nous avons besoin non seulement du mot précédent, mais aussi du mot suivant. Le modèle LSTM bidirectionnel (BiLSTM) conserve deux états distincts pour les entrées avant et arrière générées par deux LSTM différents, le premier LSTM est une séquence régulière qui commence au début de la phrase, tandis que dans le

second LSTM, la séquence d'entrée est alimentée dans l'ordre inverse. L'idée derrière le réseau bidirectionnel est de capturer des informations sur les entrées environnantes. Il apprend généralement plus rapidement qu'une approche unidirectionnelle bien que cela dépende de la tâche.

l'utilisation du LSTM bidirectionnel (BiLSTM) présente l'avantage d'obtenir des annotations de mots qui résument les informations dans les deux sens. Enfin, la raison pour laquelle une couches de BiLSTM est utilisée est de permettre au modèle d'apprendre des fonctionnalités plus abstraites.

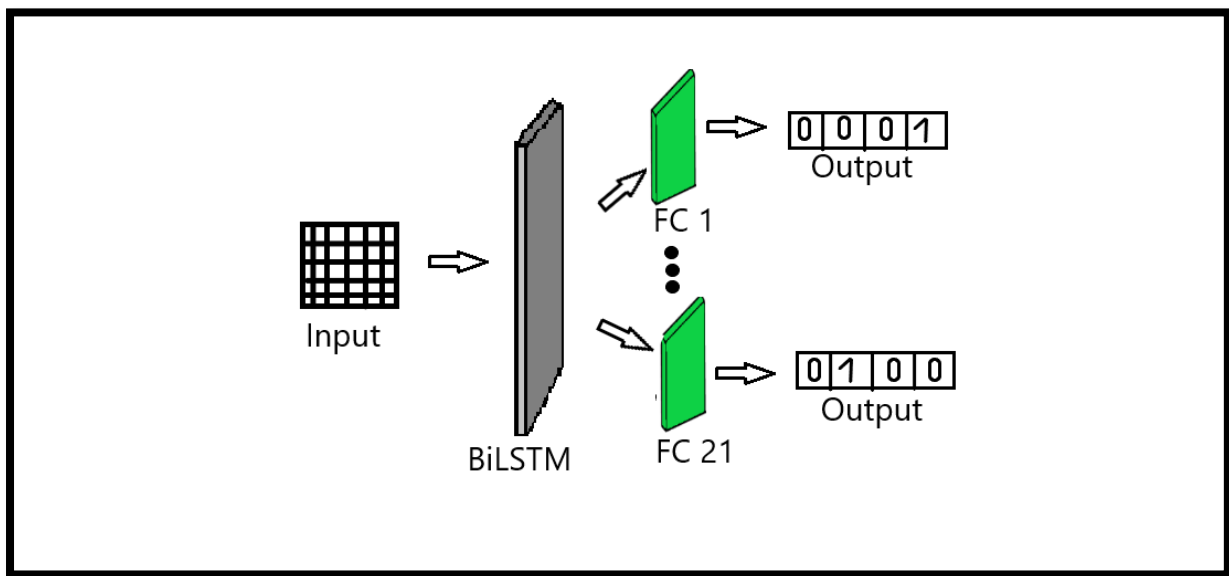


Figure 20 Architecture de notre modèle BiLSTM

5.3. Modèle biLSTM_CNN (BI-DIRECTIONAL LONG SHORT TERM MEMORY _CONVOLUTIONAL NEURAL NETWORKS)

Nous avons combiné les deux modèles le biLSTM avec le CNN, le schéma ci-dessous explique notre architecture proposée. (plus d edetaille)

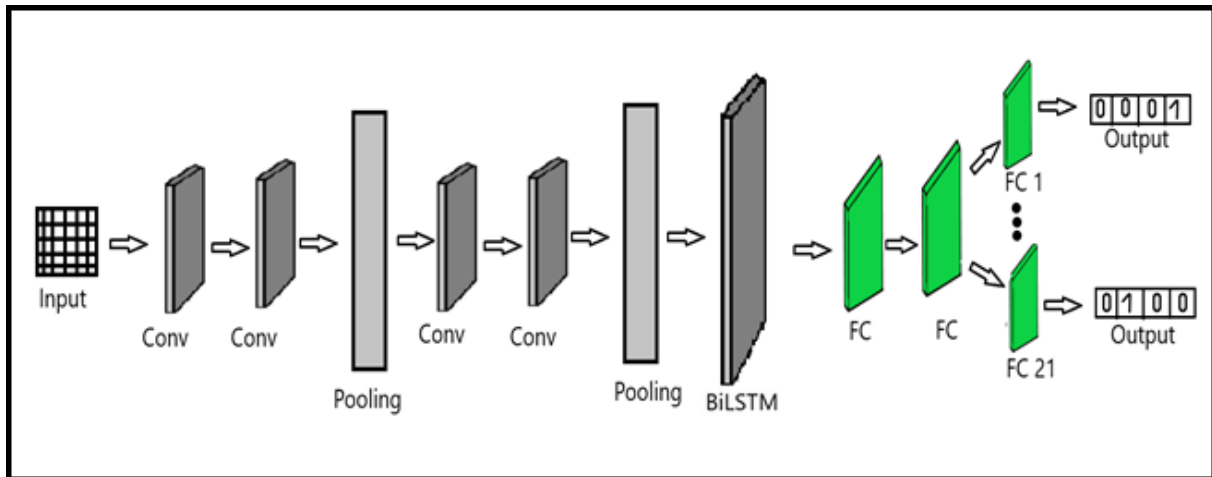


Figure 21 Architecture de notre modèle BiLSTM_CNN

6. Test (Prediction) :

Nos trois modèles contiennent une couche finale Fully connected activée par une fonction logistique nommée « softmax » ou plus communément connue sous le nom de « fonction exponentielle normalisée ».

Cette fonction attribue des probabilités décimales pour chaque réponse, dont la somme totale est de 1. Celle-ci suggère également que toute publication possède une seule réponse de chaque question.

Nos trois modèles nous donnent pour chaque message 21 sorties (outputs), chaque output est une réponse des questions de BDI. En Outre :

1 sujet $\longrightarrow$ n messages et n messages = n (21 outputs)

Tant dis qu'il faut uniquement 21 réponses seulement pour chaque sujet, nous avons utilisé deux méthodes statistiques pour générer 4 exécutions (runs). Les deux premières exécutions on été appliquées au modèle CNN et les deux autres au modèle biLSTM.

La première méthode statistique nommée max consiste à calculer pour une question la fréquence de chaque réponse générée pour choisir la réponse avec la plus grande occurrence en outre celle qui apparait le plus, c'est ce qui la rend la réponse plus efficace.

La deuxième méthode appelée suite est formé de façon proche à celle du max elle consiste à calculer pour une question, la fréquence de chaque réponse pour élire la séquence de fréquence la plus longue d'une même réponse pour toutes les réponses générées.

Exemple : Après avoir fait la prédiction de chaque publication on obtient le tableau ci-dessous en utilisant l'un de nos trois modèles nous obtenons des réponses comme celle-ci pour une seule personne.

Conception d'un modèle de mesure du degré de dépression à partir des réseaux sociaux

TEXT	Sadness	Pessimism	Past Failure	Loss of Pleasure	Guilty Feelings
the only way this couldve been a real spellcheck mistake is if they say that word a lot my sp	1	3	3	2	
youtube	1	3	0	1	
yes and idk if this will help but neil laybourn had a similar experience heres a vidhttpswww	1	3	0	1	
i tried a puff of a friends ecig and it felt greattoo great it was like a relaxing form of caffeine	0	1	1	1	
friends in high school i literally talk to three people from there all of whom i met in my seni	0	1	1	1	
rhailcorporate your post is an ad	0	0	2	0	
i find dawn of hope always wrecks me when my opponents get it down was it a heavy lifter	0	0	0	0	
so i just spent about 40 minutes going through a bunch of those studies and im not convince	0	0	0	0	
i am friends with highschool kids that hung out at the pool i lifeguarded at 6 years ago	0	0	0	0	
cant wait i have the pop art version which is my favorite right now but im super stoked for tl	1	1	2	1	
hes not playing as a true winger though its more like neymar plays at barcelona a mix betwe	0	0	0	1	
no their life does not depend on voting d their life depends on voting intelligently if theres	1	1	2	1	
the government it doesnt matter where your money goes after they steal it from you if i dic	0	0	0	1	
wait wait wait you want money	2	2	1	2	
its not an animal or a rock so what now	2	2	1	2	
waluigi	1	2	2	1	
i tried to use their three step system before and it was surprisingly harsh it did nothing for r	0	1	1	1	
to be fair i was mainly asking why he didnt use it when fighting rahvinlanfear so in those cas	1	0	1	1	
	0	0	0	0	
httpwww	0	0	0	0	
bingothe allies needed each other to win the war	0	0	0	1	
i know how hard it is to be productive or even get out of bed when youre feeling so stuck yc	1	3	3	2	
it makes absolute sense that she wont let you have people over because its not actually yo	0	1	1	1	
read the bell curve this book shows some differences in the iq of some ethnic groups	0	0	0	1	
woohoo you go that is super exciting all of your hard work has paid off	0	1	1	1	
any time i see people talking about him the say fuck you shima any time i ask who it is i just	1	1	2	1	

Tableau 5 Exemple de prédiction de chaque publication d'une personne

- La première méthode statistique nommée max consiste à calculer pour une question par exemple « sadness » la fréquence de chaque réponse générée pour choisir la réponse avec la plus grande occurrence.

Nombre de réponses 0 = 15.

Nombre de réponses 1 = 9.

Nombre de réponses 2 = 2.

Nombre de réponses 3 = 0.

Si nous calculons le max entre ces valeurs, nous trouvons que max égal à 15, donc on va dire que pour cette personne la réponse pour la question de « sadness » est égale à 0

- La deuxième méthode procède en calculant le nombre de suites de chaque réponse on prend par exemple la question pour « sadness »

La plus grande suite de 0 = 6

La plus grande suite de 1 = 3

La plus grande suite de 2 = 2

La plus grande suite de 3 = 0

Par la suite on prend le maximum entre toutes ces valeurs ici comme par hasard encore une fois la réponse 0.

Nous avons généré un fichier TXT contenant les résultats obtenus que nous avons envoyé à Erisk, ce dernier les a évalués puis nous a retransmis le score de la précision de nos modèles qui sera détaillé dans le prochain chapitre.

7. Conclusion :

Pour conclure, nous avons parcouru les différentes phases et couches de notre modèle proposé. Plus important encore, nous avons creusé profondément dans chaque partie pour donner une meilleure compréhension des bases de tous les modèles utilisés, de cette façon tout est justifié.

Dans le chapitre suivant, nous allons décrire le processus que nous avons suivi pour mettre en œuvre ce modèle et les résultats que nous avons obtenus pour pouvoir les discuter à la fin.

Chapitre 05

Expérimentations et résultats

Chapitre 05 : Expérimentations et résultats

1. Introduction

Pour l'estimation du degré de dépression, un ensemble d'outils et de matériels s'avère indispensable afin de réaliser notre modèle.

Nous allons présenter dans ce chapitre, l'ensemble de matériels utilisés, par la suite nous aborderons différents résultats qui démontreront l'efficacité de notre approche proposée, pour qu'à la fin nous discutons les résultats obtenus.

2. Matériel

Le matériel joue un rôle très important dans toutes les activités de recherche en particulier en ce qui concerne la quantité de données pouvant être gérées pendant un temps donné et le temps nécessaire pour entraîner le réseau neuronal profond et l'évaluer. Par conséquent, nous allons indiquer quel type de machine a été utilisé pour les expériences qui ont été faites et quelles sont ses caractéristiques.

Ce projet a été limité par le temps, afin d'accélérer l'entraînement de nos réseaux de neurones profonds pour les envoyer à temps nous avons pris l'initiative d'exécuter dans deux machines différentes. La première est dotée d'un Intel(R) Core(TM) i3-2330M CPU @ 2.20GHz processeur x64 avec 8.00GB RAM, la seconde possède un Intel(R) Core(TM) i7-6500U CPU @ 2.50GHz processeur x64 avec 8.00GB RAM.

Du à la difficulté imposée par ce challenge qui est le nombre insuffisant de données d'entraînement, l'espace n'a pas été un élément important à prendre en considération pour bien mener ce projet.

3. Logiciel

Afin de développer notre modèle nous avons travaillé avec un environnement de développement intégré en anglais ***integrated development environment (IDE)*** PyCharm Community Édition 2019.3.3 qui a pour objectif d'augmenter la productivité en automatisant une partie des activités et en simplifiant les opérations. **PyCharm Community Édition 2019.3.3** a été développé par l'entreprise tchèque JetBrains, c'est un logiciel multiplateformes, nous l'avons sélectionné, car nous avons utilisé **Python 3.6** comme langage de programmation. Nous avons déduit depuis les travaux antérieurs que c'est le langage le plus approprié pour de nombreuses raisons :

- sa popularité pour le traitement du texte,

- nous n'avons pas à nous soucier de la gestion de la mémoire puisque c'est automatique,
- sa simplicité, il permet d'exprimer des équations et des formules complexes et il est facile, ce qui consomme moins de temps pour l'ensemble du processus de développement,
- ses bibliothèques diverses et riches, celles dédiées au deep learning et celles utilisées pour la gestion d'autres structures de données et autres.

4. Librairies

Nous citons ci-dessous quelques librairies utilisées dans la conception de notre modèle :

- **CSV (Comma Separated Values)** : CSV est un format ouvert très populaire d'importation et d'exportation de données tel que des feuilles de calcul. Les données d'un fichier .csv sont sous forme textuelle séparées par des virgules d'où son nom *Comma Separated Values*⁸.
- **Keras** : Keras ⁹est une «API» de réseaux neuronaux de haut niveau. Elle a été développée dans le but de permettre une expérimentation rapide, prend en charge à la fois les réseaux basés sur la convolution et les réseaux récurrents (ainsi que les combinaisons des deux), et fonctionne de manière transparente sur les périphériques «CPU» et «GPU» [2].
- **Pandas** : pandas¹⁰ est une bibliothèque basée sur Python fournissant des structures de données rapides, flexibles et expressives conçues pour rendre le travail avec des données «relationnelles» ou «étiquetées» à la fois facile et intuitif. Il vise à être le bloc de construction fondamental de haut niveau pour effectuer une analyse de données pratique et réelle en Python.
- **Numpy** : Numpy¹¹ est une bibliothèque open source associée au langage Python, créée spécialement pour le calcul scientifique, notamment le calcul matriciel, tout en

⁸ <https://docs.python.org/fr/3/library/csv.html>

⁹ <https://keras.io/>

¹⁰ <https://pandas.pydata.org/>

¹¹ <https://numpy.org/>

offrant de multiples fonctions permettant la création et la manipulation de matrices, vecteurs, etc.

- **TensorFlow** : est une bibliothèque de logiciels open source pour le calcul numérique de haute performance. Son architecture flexible permet un déploiement facile du calcul sur une variété de plateformes (CPUs, GPUs, TPUs), et des ordinateurs de bureau aux clusters de serveurs en passant par les périphériques mobiles. Développé à l'origine par des chercheurs et des ingénieurs de l'équipe Google Brain au sein de l'organisation IA de Google, il est livré avec un support solide pour l'apprentissage automatique et l'apprentissage en profondeur et le noyau de calcul numérique flexible est utilisé dans de nombreux autres domaines scientifiques¹².

5. Model Word2Vec :

Pour bien mener ce projet, nous avons opté pour l'utilisation du modèle word2vec de Google gratuit qu'on peut trouver sur le web.

Word2Vec google¹³ : C'est un outil qui fournit une implémentation efficace des architectures bag-of-word et skip-gram pour le calcul des représentations vectorielles des mots. Ces représentations peuvent être utilisées par la suite dans de nombreuses applications de traitement du langage naturel et pour des recherches ultérieures.

6. Data Set :

Dans cette section nous allons discuter l'ensemble de données (Dataset) exploitées dans notre recherche, ce dernier est extrait à partir des postes de Reddit des personnes dépressives et regroupées par « Subject » puis stockées dans des fichiers « XML ».

L'ensemble de données se compose de questionnaires BDI qui ont été remplis par les utilisateurs des médias sociaux avec l'historique des écrits de chaque utilisateur. Après avoir soumis le BDI, les écrits de l'utilisateur ont été extraits juste après. Ces questionnaires originaux remplis par les utilisateurs sont les données de vérité de base appelée en anglais « ground truth data » qui serviront à l'évaluation des performances des systèmes des participants au laboratoire.

¹² <https://www.tensorflow.org>

¹³ <https://code.google.com/archive/p/word2vec/>

Expérimentations et résultats

Le but principal de ce challenge est de prédire les réponses du questionnaire de Beck's Depression Inventory BDI qui évalue la présence de sentiments comme la tristesse, le pessimisme, la perte d'énergie ... il est constitué de vingt et un items de symptômes et d'attitudes, qui décrivent une manifestation comportementale spécifique de la dépression, gradués de 0 à 3 par une série de 4 énoncés reflétant le degré de gravité du symptôme (0–9 : indique une dépression minimale, 10–18 : indique une légère dépression, 19–29 : indique une dépression modérée, 30–63 : indique une dépression sévère) [Annexe n°1]

Le Dataset de e-Risk 2019 nous a été fourni cette année comme un ensemble d'apprentissage pour entraîner notre modèle. Il contient 4 sujets « Subject » de la catégorie « **dépression minimale** (minimal depression) », 4 sujets « Subject » de la catégorie « **dépression légère** (mild depression) », 4 sujets « Subject » de la catégorie « **dépression modérée** (moderate depression) », et enfin 8 sujets « Subject » de la catégorie « **dépression sévère** (severe depression) » ce qui fait un total de 20 sujets d'étude.

Le Dataset de e-Risk 2020 est un ensemble de données qui contient 23 sujets « Subject » de la catégorie « **dépression minimale** (minimal depression) », 10 sujets « Subject » de la catégorie « **dépression légère** (mild depression) », 18 sujets « Subject » de la catégorie « **dépression modérée** (moderate depression) », et 19 sujets « Subject » de la catégorie « **dépression sévère** (severe depression) ». Le nombre total de ces sujets est 70. Cet ensemble de données sera utilisé après avoir entraîné nos modèles pour les tester et les valider.

Il nous a été demandé de développer un algorithme produisant la structure suivante :

ubject2341	1	2	3	2	3	2	3	3	1	0	0	2	3	3	2	2	2	3	2	2	0
ubject2827	1	3	3	2	3	2	2	3	2	2	1	1	2	3	1	2	1	0	2	2	1
ubject9218	2	2	1	2	1	0	3	2	1	2	2	3	1	2	2	2	0	2	2	1	1
ubject9454	1	3	3	2	1	3	3	3	2	3	2	3	3	2	2	3	1	3	3	3	0
ubject1272	1	2	2	1	0	1	3	2	1	3	3	1	2	2	2	3	0	3	1	2	2
ubject2961	1	1	1	1	0	0	1	1	1	1	1	1	1	0	2	1	0	3	1	1	1
ubject2432	1	3	3	2	3	2	2	2	2	1	1	3	1	3	2	1	2	1	2	1	0
ubject436	2	1	2	3	2	0	3	3	1	2	1	3	2	2	3	2	1	2	1	3	2
ubject3707	1	3	0	1	0	0	0	1	1	2	3	0	0	2	2	3	1	1	2	1	3

Tableau 6 Structure du fichier des réponses du questionnaire de "BDI"

Chaque ligne identifie l'auteur et les réponses estimées aux questions du BDI. Les données de vérité de base (ground truth) ont le même format [96].

7. Code source :

Ici nous présentons quelques exemples du code source du modèle CNN des commentaires on était ajouté afin d'éclairer le travail.

Expérimentations et résultats

```
19
20 path_csv = 'C:/eRisk2020_T2_TRAINING_DATA/data_csv_withStopWords/'
21 os.chdir(path_csv)
22
23 print('lire le fichier csv ...')
24 data = pd.read_csv('Dataset1.csv')
25 #remplir les cellules de TEXT qui sont vides avec des NaNs
26 data['TEXT'].replace(' ', np.nan, inplace=True)
27 #supprimer les cellules NaNs
28 data.dropna(subset=['TEXT'], inplace=True)
29
30 print('chargement du model word2vec ...')
31 #charger le model word2vec
32 model = KeyedVectors.load_word2vec_format('GoogleNews-vectors-negative300.bin', binary=True)
33
34
35 #la taille maximale de nos texts
36 maxlen = 250
37
38 print('preparation des données d'entrées ....')
39 personne = data['Personne']
40 x_Text = data['TEXT'].astype(str)
41
42 #créer un tokenizer
43 t = Tokenizer()
44 t.fit_on_texts(x_Text)
```

Figure 22 Code source 1

```
45
46 #représenter chaque mot par un nombre entier
47 #En fait, ces nombres sont l'endroit où chaque mot est stocké dans l'index des mots du tokenizer
48 sequences = t.texts_to_sequences(x_Text)
49
50 #donner à toute nos séquences la même taille maxlen
51 x_train_seq = pad_sequences(sequences, maxlen=maxlen)
52
53
54
55 #la taille de notre dictionnaire d'index (exemple de dictionnaire d'index : {'name': 1, 'my': 3, 'is': 2, 'UNK': 4})
56 num_words = len(t.word_index) + 1
57
58 #créer embedding_matrix, nous utiliserons le numéro d'index des mots afin que notre modèle puisse se référer au vecteur correspondant lorsqu'il est alimenté
59 embedding_matrix = np.zeros((num_words, 300))
60
61
62 for word, i in t.word_index.items():
63     # check if the word is in the word2vec vocab
64     if word in model.wv:
65         embedding_vector = model.wv[word]
66
67         if embedding_vector is not None:
68             embedding_matrix[i] = embedding_vector
69
70
71 print(embedding_matrix.shape)
```

Figure 23 Code source 2

Expérimentations et résultats

```
CNN_Model.py × Prediction.py × fin.py × tst.py × test.py ×
63     # check if the word is in the word2vec vocab
64     if word in model.wv:
65         embedding_vector = model.wv[word]
66
67         if embedding_vector is not None:
68             embedding_matrix[i] = embedding_vector
69
70
71 print(embedding_matrix.shape)
72
73
74 print('preparation des données de sorties ....')
75 y_Sadness = data['Sadness']
76 y_Pessimism = data['Pessimism']
77 y_Past_Failure = data['Past Failure']
78 y_Loss_of_Pleasure = data['Loss of Pleasure']
79 y_Guilty_Feelings = data['Guilty Feelings']
80 y_Punishment_Feelings = data['Punishment Feelings']
81 y_Self_Dislike = data['Self-Dislike']
82 y_Self_Criticalness = data['Self-Criticalness']
83 y_Suicidal_Thoughts_or_Wishes = data['Suicidal Thoughts or Wishes']
84 y_Crying = data['Crying']
85 y_Agitation = data['Agitation']
86 y_Loss_of_Interest = data['Loss of Interest']
87 y_Indecisiveness = data['Indecisiveness']
88 y_Worthlessness = data['Worthlessness']
89 y_Loss_of_Energy = data['Loss of Energy']
```

Figure 24 Code source 3

```
76 y_Pessimism = data['Pessimism']
77 y_Past_Failure = data['Past Failure']
78 y_Loss_of_Pleasure = data['Loss of Pleasure']
79 y_Guilty_Feelings = data['Guilty Feelings']
80 y_Punishment_Feelings = data['Punishment Feelings']
81 y_Self_Dislike = data['Self-Dislike']
82 y_Self_Criticalness = data['Self-Criticalness']
83 y_Suicidal_Thoughts_or_Wishes = data['Suicidal Thoughts or Wishes']
84 y_Crying = data['Crying']
85 y_Agitation = data['Agitation']
86 y_Loss_of_Interest = data['Loss of Interest']
87 y_Indecisiveness = data['Indecisiveness']
88 y_Worthlessness = data['Worthlessness']
89 y_Loss_of_Energy = data['Loss of Energy']
90 y_Changes_in_Sleeping_Pattern = data['Changes in Sleeping Pattern']
91 y_Irritability = data['Irritability']
92 y_Changes_in_Appetite = data['Changes in Appetite']
93 y_Concentration_Difficulty = data['Concentration Difficulty']
94 y_Tiredness_or_Fatigue = data['Tiredness or Fatigue']
95 y_Loss_of_Interest_in_Sex = data['Loss of Interest in Sex']
96
97
98 # convertir les listes d'étiquettes en tableaux NumPy avant la binarisation
99 y_Sadness = np.array(y_Sadness)
100 y_Pessimism = np.array(y_Pessimism)
101 y_Past_Failure = np.array(y_Past_Failure)
```

Figure 25 Code source 4

Expérimentations et résultats

```
98 # convertir les listes d'étiquettes en tableaux NumPy avant la binarisation
99 y_Sadness = np.array(y_Sadness)
100 y_Pessimism = np.array(y_Pessimism)
101 y_Past_Failure = np.array(y_Past_Failure)
102 y_Loss_of_Pleasure = np.array(y_Loss_of_Pleasure)
103 y_Guilty_Feelings = np.array(y_Guilty_Feelings)
104 y_Punishment_Feelings = np.array(y_Punishment_Feelings)
105 y_Self_Dislike = np.array(y_Self_Dislike)
106 y_Self_Criticalness = np.array(y_Self_Criticalness)
107 y_Suicidal_Thoughts_or_Wishes = np.array(y_Suicidal_Thoughts_or_Wishes)
108 y_Crying = np.array(y_Crying)
109 y_Agitation = np.array(y_Agitation)
110 y_Loss_of_Interest = np.array(y_Loss_of_Interest)
111 y_Indecisiveness = np.array(y_Indecisiveness)
112 y_Worthlessness = np.array(y_Worthlessness)
113 y_Loss_of_Energy = np.array(y_Loss_of_Energy)
114 y_Changes_in_Sleeping_Pattern = np.array(y_Changes_in_Sleeping_Pattern)
115 y_Irritability = np.array(y_Irritability)
116 y_Changes_in_Appetite = np.array(y_Changes_in_Appetite)
117 y_Concentration_Difficulty = np.array(y_Concentration_Difficulty)
118 y_Tiredness_or_Fatigue = np.array(y_Tiredness_or_Fatigue)
119 y_Loss_of_Interest_in_Sex = np.array(y_Loss_of_Interest_in_Sex)
120
121 # binariser les deux ensembles d'étiquettes
122 SadnessLB = LabelBinarizer()
123 PessimismLB = LabelBinarizer()
124 Past_FailureLB = LabelBinarizer()
```

Figure 26 Code source 5

```
144 y_Sadness = SadnessLB.fit_transform(y_Sadness)
145 y_Pessimism = PessimismLB.fit_transform(y_Pessimism)
146 y_Past_Failure = Past_FailureLB.fit_transform(y_Past_Failure)
147 y_Loss_of_Pleasure = Loss_of_PleasureLB.fit_transform(y_Loss_of_Pleasure)
148 y_Guilty_Feelings = Guilty_FeelingsLB.fit_transform(y_Guilty_Feelings)
149 y_Punishment_Feelings = Punishment_FeelingsLB.fit_transform(y_Punishment_Feelings)
150 y_Self_Dislike = Self_DislikeLB.fit_transform(y_Self_Dislike)
151 y_Self_Criticalness = Self_CriticalnessLB.fit_transform(y_Self_Criticalness)
152 y_Suicidal_Thoughts_or_Wishes = Suicidal_Thoughts_or_WishesLB.fit_transform(y_Suicidal_Thoughts_or_Wishes)
153 y_Crying = CryingLB.fit_transform(y_Crying)
154 y_Agitation = AgitationLB.fit_transform(y_Agitation)
155 y_Loss_of_Interest = Loss_of_InterestLB.fit_transform(y_Loss_of_Interest)
156 y_Indecisiveness = IndecisivenessLB.fit_transform(y_Indecisiveness)
157 y_Worthlessness = WorthlessnessLB.fit_transform(y_Worthlessness)
158 y_Loss_of_Energy = Loss_of_EnergyLB.fit_transform(y_Loss_of_Energy)
159 y_Changes_in_Sleeping_Pattern = Changes_in_Sleeping_PatternLB.fit_transform(y_Changes_in_Sleeping_Pattern)
160 y_Irritability = IrritabilityLB.fit_transform(y_Irritability)
161 y_Changes_in_Appetite = Changes_in_AppetiteLB.fit_transform(y_Changes_in_Appetite)
162 y_Concentration_Difficulty = Concentration_DifficultyLB.fit_transform(y_Concentration_Difficulty)
163 y_Tiredness_or_Fatigue = Tiredness_or_FatigueLB.fit_transform(y_Tiredness_or_Fatigue)
164 y_Loss_of_Interest_in_Sex = Loss_of_Interest_in_SexLB.fit_transform(y_Loss_of_Interest_in_Sex)
165
166 target_cols = [y_Sadness, y_Pessimism, y_Past_Failure, y_Loss_of_Pleasure, y_Guilty_Feelings, y_Punishment_Feelings, y_Self_Dislike, y_Self_Criticalness, y_Suicidal_Thoughts_or_Wishes,
167               y_Loss_of_Interest, y_Indecisiveness, y_Worthlessness, y_Loss_of_Energy, y_Changes_in_Sleeping_Pattern, y_Irritability, y_Changes_in_Appetite,
168               y_Concentration_Difficulty, y_Tiredness_or_Fatigue, y_Loss_of_Interest_in_Sex]
169
170 print("division du dataset en testdata et testdata")
```

Figure 27 Code source 6

Expérimentations et résultats

```
169 y_concentration_difficulty, y_tiredness_or_fatigue, y_loss_of_interest_in_sex])
170 print('division du dataset en traindata et testdata ...')
171 split = train_test_split(personne, x_train_seq, y_Sadness, y_Pessimism, y_Past_Failure, y_Loss_of_Pleasure, y_Guilty_Feelings,
172 y_Punishment_Feelings, y_Self_Dislike, y_Self_Criticalness, y_Suicidal_Thoughts_or_Wishes,
173 y_Crying, y_Agitation, y_Loss_of_Interest, y_Indecisiveness, y_Worthlessness, y_Loss_of_Energy,
174 y_Changes_in_Sleeping_Pattern, y_Irritability, y_Changes_in_Appetite, y_Concentration_Difficulty,
175 y_Tiredness_or_Fatigue, y_Loss_of_Interest_in_Sex, test_size=0.2, random_state=42)
176
177 (train_personne, test_personne, train_x_Text, test_x_Text, train_y_Sadness, test_y_Sadness, train_y_Pessimism, test_y_Pessimism, train_y_Past_Failure, test_y_Past_
178 train_y_Punishment_Feelings, test_y_Punishment_Feelings, train_y_Self_Dislike, test_y_Self_Dislike, train_y_Self_Criticalness, test_y_Self_Criticalness
179 train_y_Crying, test_y_Crying, train_y_Agitation, test_y_Agitation, train_y_Loss_of_Interest, test_y_Loss_of_Interest, train_y_Indecisiveness, test_y_In
180 train_y_Changes_in_Sleeping_Pattern, test_y_Changes_in_Sleeping_Pattern, train_y_Irritability, test_y_Irritability, tain_y_Changes_in_Appetite, test_y_
181 train_y_Tiredness_or_Fatigue, test_y_Tiredness_or_Fatigue, train_y_Loss_of_Interest_in_Sex, test_y_Loss_of_Interest_in_Sex) = split
182
183 y_train = [train_y_Sadness, train_y_Pessimism, train_y_Past_Failure, train_y_Loss_of_Pleasure, train_y_Guilty_Feelings,
184 train_y_Punishment_Feelings, train_y_Self_Dislike, train_y_Self_Criticalness, train_y_Suicidal_Thoughts_or_Wishes,
185 train_y_Crying, train_y_Agitation, train_y_Loss_of_Interest, train_y_Indecisiveness, train_y_Worthlessness, train_y_Loss_of_Energy,
186 train_y_Changes_in_Sleeping_Pattern, train_y_Irritability, tain_y_Changes_in_Appetite, train_y_Concentration_Difficulty,
187 train_y_Tiredness_or_Fatigue, train_y_Loss_of_Interest_in_Sex]
188
189 y_test = [test_y_Sadness, test_y_Pessimism, test_y_Past_Failure, test_y_Loss_of_Pleasure, test_y_Guilty_Feelings
190 test_y_Punishment_Feelings, test_y_Self_Dislike, test_y_Self_Criticalness, test_y_Suicidal_Thoughts_or_Wishes
191 test_y_Crying, test_y_Agitation, test_y_Loss_of_Interest, test_y_Indecisiveness, test_y_Worthlessness, test_y_Loss_of_Energy
192 test_y_Changes_in_Sleeping_Pattern, test_y_Irritability, test_y_Changes_in_Appetite, test_y_Concentration_Difficulty
193 test_y_Tiredness_or_Fatigue, test_y_Loss_of_Interest_in_Sex]
194
```

Figure 28 Code source 7

```
198
199 print('preparation du model')
200 text_input = Input(shape=(250,))
201 #e = Embedding(num_words, 300, weights=[embedding_matrix], input_length=250, trainable=True)(text_input)
202 e = Embedding(num_words, 300, weights=[embedding_matrix], input_length=250, trainable=False)(text_input)
203
204 layers = Conv1D(filters=64, kernel_size=3, activation='relu')(e)
205 layers = Conv1D(filters=32, kernel_size=3, activation='relu')(layers)
206 layers = MaxPooling1D()(layers)
207 layers = Conv1D(filters=32, kernel_size=3, activation='relu')(layers)
208 layers = Conv1D(filters=32, kernel_size=3, activation='relu')(layers)
209 layers = Flatten()(layers)
210 layers = Dense(500, activation='relu')(layers)
211 layers = Dense(200, activation='relu')(layers)
212
213
214 output1 = Dense(3, activation='softmax')(layers)
215 output2 = Dense(4, activation='softmax')(layers)
216 output3 = Dense(4, activation='softmax')(layers)
217 output4 = Dense(4, activation='softmax')(layers)
218 output5 = Dense(4, activation='softmax')(layers)
219 output6 = Dense(4, activation='softmax')(layers)
220 output7 = Dense(4, activation='softmax')(layers)
221 output8 = Dense(4, activation='softmax')(layers)
222 output9 = Dense(3, activation='softmax')(layers)
223 output10 = Dense(4, activation='softmax')(layers)
```

Figure 29 Code source 8

Expérimentations et résultats

```
CNN_Model.py x Prediction.py x fin.py x tst.py x test.py x
212
213
214 output1 = Dense(3, activation='softmax')(layers)
215 output2 = Dense(4, activation='softmax')(layers)
216 output3 = Dense(4, activation='softmax')(layers)
217 output4 = Dense(4, activation='softmax')(layers)
218 output5 = Dense(4, activation='softmax')(layers)
219 output6 = Dense(4, activation='softmax')(layers)
220 output7 = Dense(4, activation='softmax')(layers)
221 output8 = Dense(4, activation='softmax')(layers)
222 output9 = Dense(3, activation='softmax')(layers)
223 output10 = Dense(4, activation='softmax')(layers)
224 output11 = Dense(4, activation='softmax')(layers)
225 output12 = Dense(4, activation='softmax')(layers)
226 output13 = Dense(4, activation='softmax')(layers)
227 output14 = Dense(4, activation='softmax')(layers)
228 output15 = Dense(4, activation='softmax')(layers)
229 output16 = Dense(4, activation='softmax')(layers)
230 output17 = Dense(3, activation='softmax')(layers)
231 output18 = Dense(4, activation='softmax')(layers)
232 output19 = Dense(4, activation='softmax')(layers)
233 output20 = Dense(4, activation='softmax')(layers)
234 output21 = Dense(4, activation='softmax')(layers)
235
236 model_cnn=Model(inputs=text_input,outputs=[output1,output2,output3,output4,output5,output6,output7,output8,output9,output10,output11,output12,output13,output14,output15,output16,output17,output18,output19,output20,output21])
237
238 model_cnn.compile(loss='binary_crossentropy',
```

Figure 30 Code source 9

```
225 output12 = Dense(4, activation='softmax')(layers)
226 output13 = Dense(4, activation='softmax')(layers)
227 output14 = Dense(4, activation='softmax')(layers)
228 output15 = Dense(4, activation='softmax')(layers)
229 output16 = Dense(4, activation='softmax')(layers)
230 output17 = Dense(3, activation='softmax')(layers)
231 output18 = Dense(4, activation='softmax')(layers)
232 output19 = Dense(4, activation='softmax')(layers)
233 output20 = Dense(4, activation='softmax')(layers)
234 output21 = Dense(4, activation='softmax')(layers)
235
236 model_cnn=Model(inputs=text_input,outputs=[output1,output2,output3,output4,output5,output6,output7,output8,output9,output10,output11,output12,output13,output14,output15,output16,output17,output18,output19,output20,output21])
237
238 model_cnn.compile(loss='binary_crossentropy',
239                   optimizer='adam',
240                   metrics=['accuracy'])
241
242 model_cnn.summary()
243 checkpoint = ModelCheckpoint('cnn_depression.hdf5 ', monitor='val_acc', verbose=1, save_best_only=True, mode='max')
244
245 model_cnn.fit(train_x_Text,y_train, batch_size=50, epochs=100,validation_data=(test_x_Text,y_test))
246 model_cnn.save("model3.h5")
```

Figure 31 Code source 10

8. Reddit :

C'est une plate-forme open source où les membres de la communauté (red-ditors) peuvent soumettre du contenu (articles, commentaires ou liens directs), soumettre des votes, et les entrées de contenu sont organisées par domaines d'intérêt (subreddits). Reddit a une grande communauté de membres et de nombreux membres ont une grande histoire de soumissions précédentes (couvrant plusieurs années). Il contient également un contenu substantiel sur différentes conditions médicales, telles que l'anorexie ou la dépression.

Les conditions générales de Reddit permettent d'utiliser son contenu à des fins de recherche.

Reddit remplit tous les critères de sélection cités par Losada et tous [96], par conséquent, ils l'ont utilisé pour créer la collection de tests de dépression.

9. Mesure de performance

La tâche utilise une variété de mesures d'évaluation pour mesurer le succès des algorithmes. Losada et tous [96]. Les définissent comme suit :

Le taux de réussite « Hit Rate (HR) » : HR détermine à quelle fréquence la prédiction était correcte, par rapport au questionnaire réel et donne le ratio. Par exemple, une prédiction où 5 des 21 questions du BDI correct ont obtenu une valeur **HR** de 5/21.

Le taux de réussite moyen « Average Hit Rate (AHR) » : Le AHR est HR, mais en moyenne pour tous les utilisateurs.

« Closeness Rate CR » : CR considère l'échelle ordinale sous-jacente aux questions du BDI. Pour chaque question, une différence absolue (da) entre la réponse réelle et la réponse prévue. Un système qui est plus éloigné de la réponse qu'un deuxième système devrait être pénalisé pour cette plus grande distance. Pour cela, la mesure est construite comme ceci :

$$CR = \frac{(mda - da)}{mda}$$

Ici mda représente la différence absolue maximale qui est le nombre de réponses possibles moins un.

Average Closeness Rate (ACR) : ACR est CR, mais en moyenne sur tous les utilisateurs. Cependant, les questions # 16 et # 18 ont sept possibilités de réponse, où pour les réponses 1 à 3, deux options possibles (a et b) sont disponibles. Cependant, ces options ont été considérées comme égales, car elles représentent le même niveau de dépression.

Difference between overall depression levels (DODL) : cette mesure ne prend pas en compte les prédictions correctes du système au niveau des questions, mais donne le niveau de dépression global basé sur la somme de toutes les réponses pour le BDI réel et généré par le système. De plus, la différence absolue (da globale) entre le score de dépression réel et le score de dépression prévu est calculée.

Expérimentations et résultats

Un niveau de dépression est un entier compris entre 0 et 63. Ces nombres sont dérivés de l'addition des nombres de réponses du BDI. Par exemple, en considérant la question n ° 1 [annexe A], si la réponse était l'option 1, le niveau de dépression entier est augmenté de 1. De cette façon, les quatre catégories suivantes sont associées aux niveaux de dépression respectifs:

- Dépression minimale (niveaux de dépression 0-9)
- Dépression légère (niveaux de dépression 10-18)
- Dépression modérée (niveaux de dépression 19-29)
- Dépression sévère (niveaux de dépression 30-63)

Ces niveaux sont largement acceptés dans la littérature psychologique [101].

La mesure DODL est finalement normalisée dans l'intervalle [0,1] de la manière suivante :

$$\text{DODL} = \frac{63 - da_{globale}}{63}$$

Average DODL (ADODL) : ADODL est DODL, mais en moyenne pour tous les utilisateurs.

Depression Category Hit Rate (DCHR) : DCHR calcule la fraction de cas, dans laquelle le système où est généré BDI, a conduit à la même catégorie de dépression obtenue à partir du questionnaire des réels auteurs.

10. Résultat et évaluations des performances

Cette section présente les résultats de nos expériences et les compare aux résultats des exécutions soumises pour la tâche à CLEF eRisk2020.

Cette année, 6 équipes au total ont soumis 17 essais différents à la tâche 2 eRisk . Notre équipe a soumis 4 exécutions à cette tâche. Les diagrammes ci-dessous montrent nos résultats comparés aux meilleurs résultats pour chaque métrique. Les résultats de tous les participants peuvent être trouvés dans [96]. Nous avons observé qu'aucune équipe n'a été en mesure d'obtenir les meilleurs résultats pour chacune des évaluations des quatre métriques.

En comparant uniquement nos exécutions, les exécutions 2 et 4 utilisant la technique statistique suites de mêmes réponses expliquée dans le chapitre précédent n'ont pas bien fonctionné, les résultats obtenus pour ses deux exécutions n'ont pas été concluants. En revanche la première exécution utilisant CNN avec la technique max de chaque réponse a obtenu les meilleurs résultats en AHR avec **34,97%** des réponses correctes, et le meilleur dans la mesure DCHR qui est capable de prédire la bonne catégorie de gravité de dépression pour

Expérimentations et résultats

25,71% des utilisateurs. Contrairement aux trois exécutions, l'exécution 3 a atteint un record d'ACR égale à **67,78%** et d'ADODL égale à **79,30%** les records les plus majoritaires obtenus de nos exécutions .

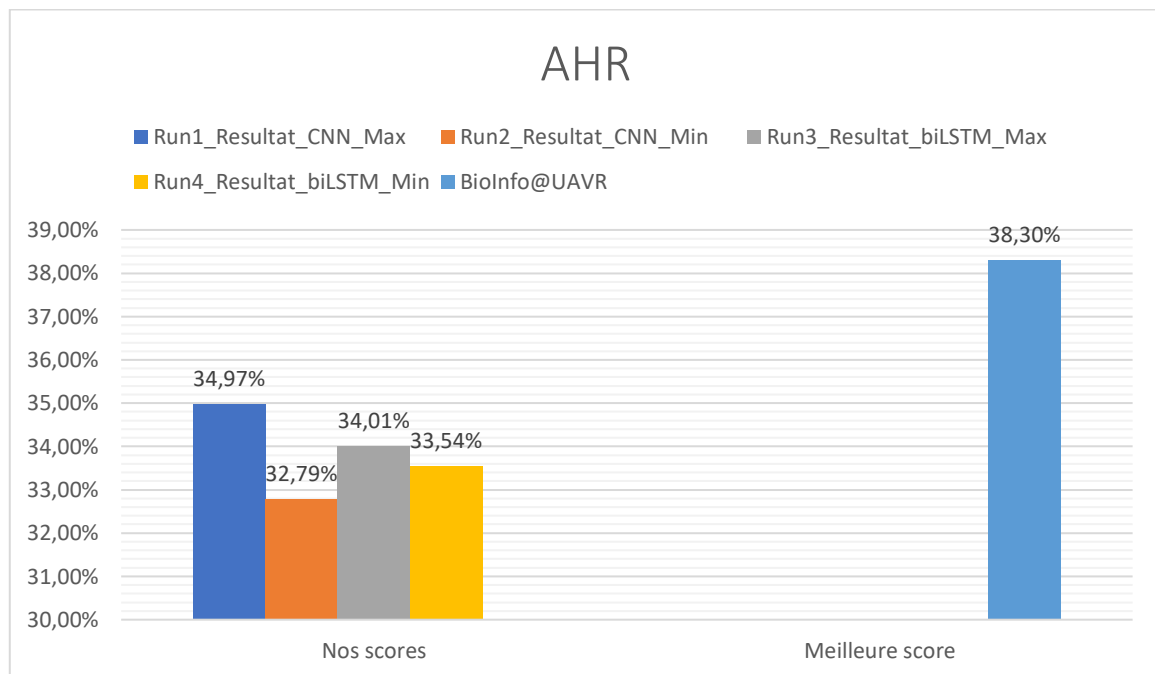


Diagramme 1 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure AHR au challenge eRisk2020

Expérimentations et résultats

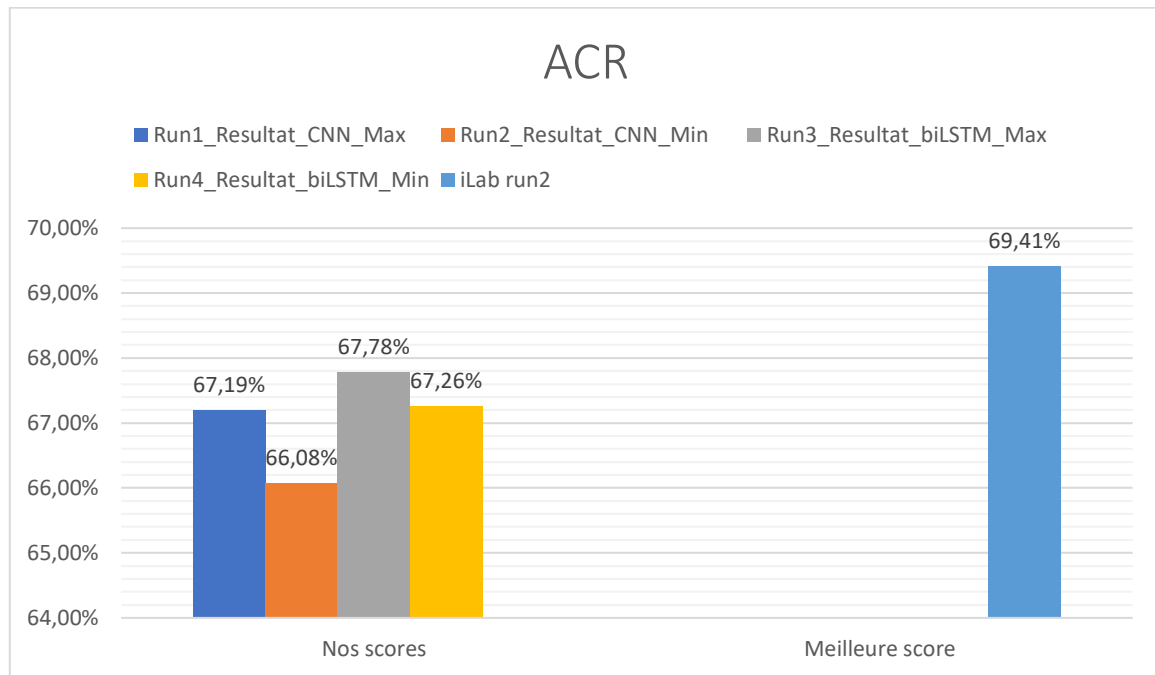


Diagramme 2 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure ACR au challenge eRisk2020

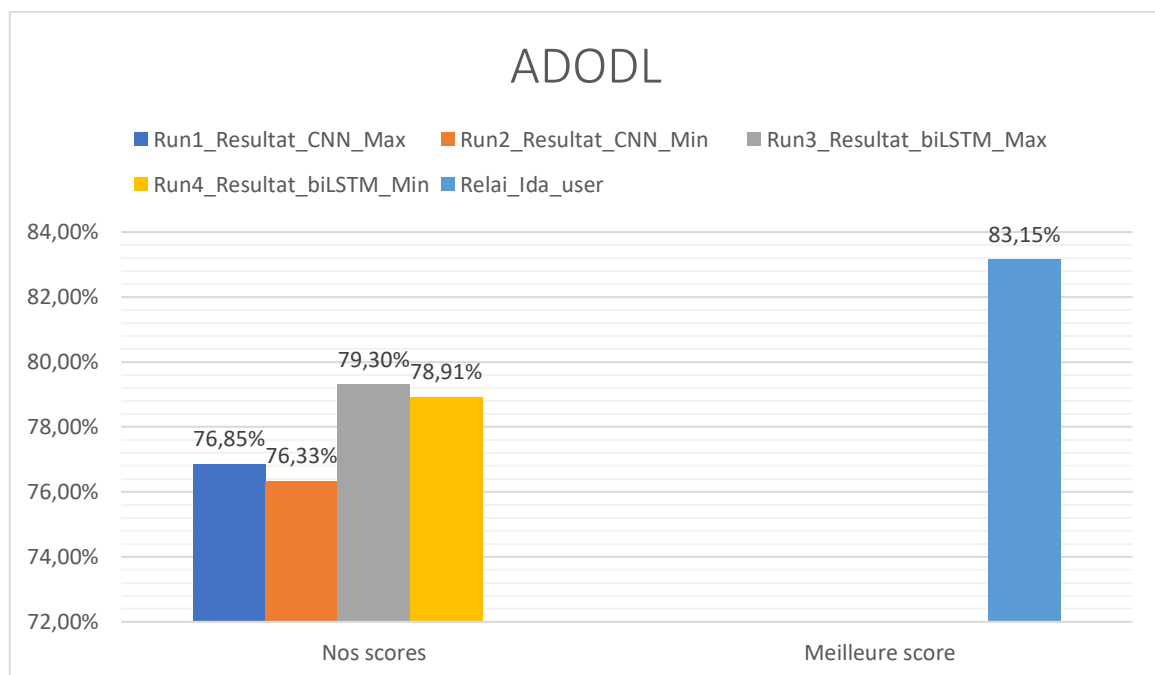


Diagramme 3 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure ADODL au challenge eRisk2020

Expérimentations et résultats

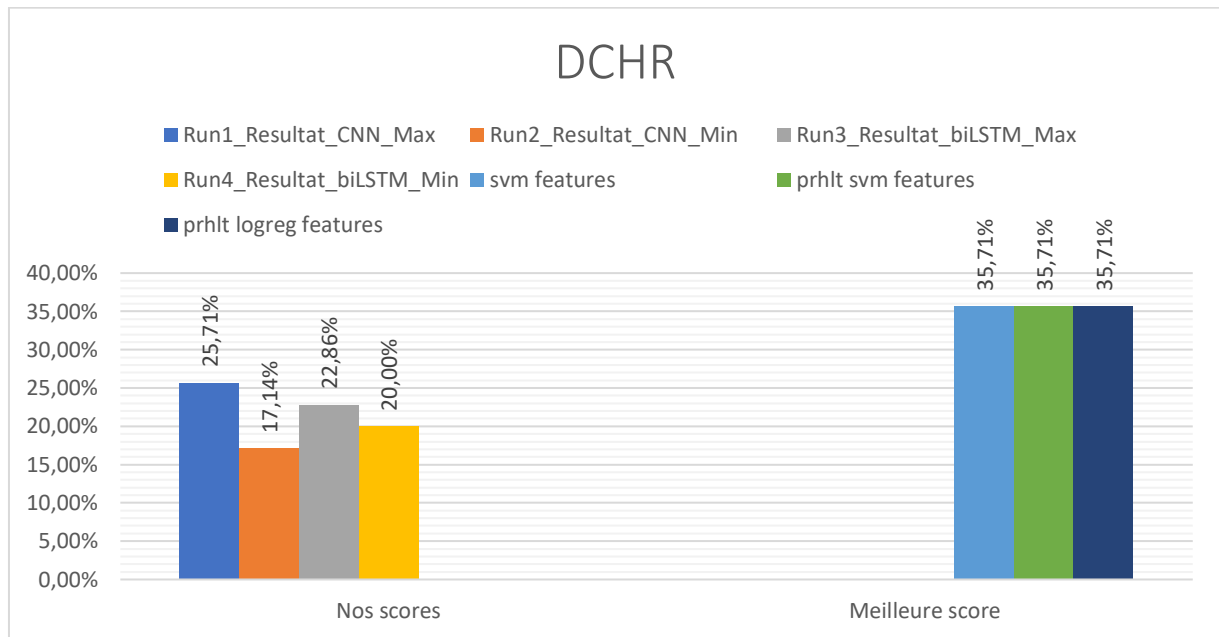


Diagramme 4 Comparaison de nos résultats avec les meilleurs résultats obtenus pour la mesure DCHR au challenge eRisk2020

Suite à la comparaison et aux diagrammes nous avons conclu que l'utilisation de la première méthode qui consiste à calculer le max de chaque réponse et qu'on finit par trancher avec la réponse qui réapparaît à chaque fois, est bien plus performante que celle des suites de mêmes réponses. Aussi il s'avère que le réseau de neurones convolutif dans le modèle qu'on a proposé pour prédire la sévérité de dépression donne légèrement de meilleurs résultats comparés au biLSTM.

Répondre au questionnaire BDI avec un ensemble de données composé de 20 sujets d'entraînement seulement s'est avéré être un défi de taille. Le tableau ci-dessous résume les résultats de tous les participants à la tâche 2 de Erisk 2020. Globalement les résultats de cette tâche sont assez homogènes avec peu de variations d'une équipe à l'autre. Cela peut être vu comme une indication à la difficulté de la tâche.

Executions	AHR	ACR	ADODL	DCHR
BioInfo@UAVR	38.30%	69.21%	76.01%	30.00%
iLab run1	36.73%	68.68%	81.07%	27.14%
iLab run2	37.07%	69.41%	81.70%	27.14%

Expérimentations et résultats

iLab run3	35.99%	69.14%	82.93%	34.29%
prhlt logreg features	34.01%	67.07%	80.05%	35.71%
prhlt svm use	36.94%	69.02%	81.72%	31.43%
prhlt svm features	34.56%	67.44%	80.63%	35.71%
svm features	34.56%	67.44%	80.63%	35.71%
relai context paral user	36.80%	68.37%	80.84%	22.86%
relai context sim answer	21.16%	55.40%	73.76%	27.14%
relai lda answer	28.50%	60.79%	79.07%	30.00%
relai lda user	36.39%	68.32%	83.15%	34.29%
relai sylo user	37.28%	68.37%	80.70%	20.00%
Run1 resultat CNN Methode max	34.97%	67.19%	76.85%	25.71%
Run2 resultat CNN Methode suite	32.79%	66.08%	76.33%	17.14%
Run3 resultat BILSTM Methode max	34.01%	67.78%	79.30%	22.86%
Run4 resultat BILSTM Methode suit	33.54%	67.26%	78.91%	20.00%

Tableau 7 montrant la performance des résultats eRisk2020

Nous avons testé nos modèles sur le dataset de eRisk 2019, les résultats sont :

Les exécutions	AHR	ACR	ADODL	DCHR
1.resultat_CNN_Methode_Max2019	40.24%	67.30%	72.38%	30%

Expérimentations et résultats

2.resultat_BILSTM_Methode_Max2019	42.38%	70.32%	78.10%	25%
3.resultat_CNN_Methode_suite2019	32.86%	65.32%	69.44%	15%
4.resultat_BILSTM_Methode_suite2019	40%	68.89%	78.73%	25%

Tableau 8 récapitulant Nos Resultats eRisk2019

Les exécutions	AHR	ACR	ADODL	DCHR
BioInfo@UAVR	34.05%	66.43%	77.70%	25.00%
BiTeM	32.14%	62.62%	72.62%	25.00%
CAMHGPTnearestunsupervised	23.81%	57.06%	81.03%	45.00%
CAMHGPTsupervised.181features.58hr	35.47%	68.33%	75.63%	20.00%
CAMHGPTsupervised.769features.55hr	36.43%	67.22%	72.30%	20.00%
CAMHGPTsupervised.949features.75hr	36.91%	69.13%	75.63%	15.00%
CAMHLIWCsupervisedSVM	35.95%	66.59%	75.48%	25.00%
Fazl	22.38%	56.27%	72.78%	5.00%
Illinois	22.62%	56.19%	66.35%	40.00%
ISIKolmultiSimilarity-5000-Dtac-Qtac	29.76%	57.94%	74.13%	25.00%
ISIKol-bm25-1.2-0.75-5000-Dtac-Qtac	29.76%	57.06%	72.78%	25.00%
ISIKol-lm-d-1.0-5000-Dtac-Qtac	30.00%	57.94%	73.02%	15.00%
Kimberly	38.33%	64.44%	66.19%	20.00%

Expérimentations et résultats

UNSLA	37.38%	67.94%	72.86%	30.00%
UNSLB	36.93%	70.16%	76.83%	30.00%
UNSLC	41.43%	69.13%	78.02%	40.00%
UNSLD	38.10%	67.22%	78.02%	30.00%
UNSL E	40.71%	71.27%	80.48%	35.00%
1. resultat_CNN_Methode_Max2019	40.24%	67.30%	72.38%	30%
2. resultat_BiLSTM_Methode_Max2019	42.38%	70.32%	78.10%	25%
3. resultat_CNN_Methode_suite2019	32.86%	65.32%	69.44%	15%
4. resultat_BiLSTM_Methode_suite2019	40%	68.89%	78.73%	25%

Tableau 9 La performance des résultats eRisk 2019 associés à nos résultats [5].

Ici nous remarquons qu'on a eu le meilleur score en ACR **42.38%** avec la méthode max du biLSTM, et des résultats très proches pour les autres mesures. Ceci nous mène à croire que combiner le modèle CNN avec celui du biLSTM devrait accroître le processus d'extraction des fonctionnalités et améliorer les performances du modèle pour prédire de meilleurs résultats.

Nous avons alors combiné le modèle biLSTM au CNN avec la méthode statistique max sur le dataset eRisk2020 les résultats obtenus sont plutôt concluants nous avons obtenu la deuxième place dans la métrique AHR avec un score de **37,41%** et des résultats proches pour le reste avec un score de **67,17%** d'ACR et **73,65%** d'ADODL et pour finir **18,57%** DCHR .

Tant dis qu'avec la méthode statistique suite les résultats sont moins satisfaisants avec un score : ACR **34,83%** , AHR **66,39%** , ADODL **74,24%** , DCHR **14,29%**.

11.Conclusion

Pour conclure, ce chapitre d'expérimentation a été la partie la plus intéressante de toute cette étude, car nous avons évalué les différents résultats obtenus nous sommes allés même jusqu'à mesurer la performance de nos modèles avec le challenge de l'an passé. Nous avons aussi

Expérimentations et résultats

exprimé les différents outils utilisés. Cela étant dit, c'était la tâche la plus longue à faire pour pouvoir enfin conclure et passer à la conclusion générale.

Conclusion générale

Conclusion générale

L'objectif de notre étude était d'atténuer les risques sanitaires associés aux troubles majeurs dépressifs, en raison de sa gravité à affecter le mental tout comme le physique. Cette maladie affecte la vie quotidienne de la personne et l'empoisonne. Cette étude a été menée pour éviter les tragédies qui pourraient survenir à la suite d'un tel désordre comme le suicide. Pour cela, nous avons voulu concevoir et proposer une nouvelle façon d'estimer le degré de dépression à l'aide d'une technologie avancée de l'intelligence artificielle appelée le deep learning.

Pour y parvenir, nous avons examiné autant d'études que possible sur la question. Tout d'abord, nous avons mené une recherche sur l'analyse des sentiments et passé en revue une grande variété de travaux liés, mais sans s'y limiter, à la santé mentale. Nous avons ensuite étudié les différents modèles existants de l'apprentissage en profondeur. Ensuite, nous avons passé en revue les différentes approches qui abordaient la même problématique que la nôtre, plus précisément les travaux liés à la mesure du niveau de dépression. Cela nous a également permis de remarquer que des technologies telles que l'apprentissage en profondeur n'étaient pas exploitées pour résoudre ces problèmes, et aussi nous avons rencontré quelques difficultés à trouver des articles dans ce domaine ce qui montre que ce n'est que récemment que les gens commencent à s'intéresser à ces problèmes. À la fin de ce processus, nous avons une compréhension suffisante de toutes ces technologies et de l'état de l'art actuel pour être en mesure de développer et de proposer notre propre modèle pour atteindre notre objectif principal.

Nous avons proposé trois modèles en profitant d'un modèle existant qui était auparavant utilisé strictement pour l'analyse des sentiments, et nous avons essayé de l'inclure dans nos modèles d'estimation du niveau de dépression avec nos améliorations. Nos modèles exploitent principalement trois types de réseaux neuronaux profonds, un réseau neuronal convolutif « CNN » et un réseau de neurones récurrent bidirectionnel à longue mémoire à court terme « biLSTM », et une combinaison des deux « biLSTM_CNN ». Pour l'entraînement de ses derniers, nous avons utilisé le dataset que CLEF eRisk nous a généreusement fournie. Pour ce faire, nous avons fait recours à des méthodes statistiques spécifiques. L'approche qu'on a utilisée pour nos prédictions est non supervisée. Grâce à notre participation au challenge eRisk 2020, nous avons eu la chance de nous faire évaluer nos résultats par le système d'évaluation de CLEF eRisk2020. Nous avons même demandé une évaluation de nos modèles par rapport au dataset eRisk2019. Dans l'ensemble, les résultats sont très avantageux dont nous avons obtenu le second meilleur score dans la première métrique d'évaluation « ACR ».

Nos principales contributions à ce travail reposent sur les trois points suivants :

Conclusion générale

Nous avons combiné le modèle CNN avec le modèle biLSTM afin d'améliorer le processus d'extraction de caractéristiques et d'améliorer les performances du modèle.

Nous avons utilisé nos propres méthodes statistiques pour la prédiction.

Nous avons analysé à travers différentes expériences les performances de trois modèles d'apprentissage profond afin de fournir plus de perspectives et

Nous considérons cette étude comme une première étape vers le dévoilement de nouvelles façons d'estimer le degré de la dépression et d'autres maladies mentales, non seulement sur Reddit, mais également dans différents types de plates-formes et d'applications pouvant être utilisées au quotidien.

Comme perspectives d'avenir, premièrement, nous suggérons qu'un plus complexe modèle soit développé à partir du nôtre, ce nouveau devrait pouvoir prendre en considération non seulement les publications Reddit, mais aussi leur timing et d'autres métadonnées ainsi. Un tel modèle devrait pouvoir traiter plus que les publications Reddit d'un utilisateur donné, mais aussi tous ses détails, cela comprend: le pays, la photo de profil, le nombre de followers, l'âge... etc. et tous les détails possibles qui peuvent être récupérés d'un compte Reddit.

Deuxièmement, nous croyons que notre modèle peut être utilisé avec différentes maladies mentales, pas seulement la dépression. La seule partie manquante pour le prouver est des ensembles de données dédiés spécifiquement à divers troubles mentaux. Nous pensons également que notre modèle peut être utilisé sur n'importe quel corpus tant qu'il contient des déclarations auto-exprimées et qu'il fournira les mêmes résultats.

Enfin, notre perspective la plus optimiste et futuriste est que le modèle soit amélioré de façon à ce qu'au lieu de prédire pour chaque publications 21 sorties, on parcourt tout le sujet une seule fois entièrement pour en extraire les réponses nécessaires une seule fois, en outre 21 sorties pour répondre au BDI.

Pour conclure, ce projet a été une expérience d'apprentissage assez enrichissante, car il nous a fait réaliser à quel point la dépression peut être une maladie grave qu'il faudrait vraiment prendre au sérieux. Nous avons appris de nouvelles technologies inconnues et le domaine de l'analyse des sentiments et enfin cette expérience nous a permis d'acquérir de nouvelles compétences et de maîtriser celles existantes, cette expérience nous a permis de murir et gagné de l'expérience nous avons publié notre premier paper [annexe B] durant notre participation au challenge CLEF eRisk 2020 .

Annexe

[Annexe A] : Beck's Depression Inventory

1. Sadness

- 0. I do not feel sad.
- 1. I feel sad much of the time.
- 2. I am sad all the time.
- 3. I am so sad or unhappy that I can't stand it.

2. Pessimism

- 0. I am not discouraged about my future.
- 1. I feel more discouraged about my future than I used to.
- 2. I do not expect things to work out for me.
- 3. I feel my future is hopeless and will only get worse.

3. Past Failure

- 0. I do not feel like a failure.
- 1. I have failed more than I should have.
- 2. As I look back, I see a lot of failures.
- 3. I feel I am a total failure as a person.

4. Loss of Pleasure

- 0. I get as much pleasure as I ever did from the things I enjoy.
- 1. I don't enjoy things as much as I used to.
- 2. I get very little pleasure from the things I used to enjoy.
- 3. I can't get any pleasure from the things I used to enjoy.

5. Guilty Feelings

- 0. I don't feel particularly guilty.
- 1. I feel guilty over many things I have done or should have done.
- 2. I feel quite guilty most of the time.
- 3. I feel guilty all of the time.

6. Punishment Feelings

- 0. I don't feel I am being punished.
- 1. I feel I may be punished.
- 2. I expect to be punished.
- 3. I feel I am being punished.

7. Self-Dislike

- 0. I feel the same about myself as ever.
- 1. I have lost confidence in myself.
- 2. I am disappointed in myself.
- 3. I dislike myself.

8. Self-Criticalness

- 0. I don't criticize or blame myself more than usual.
- 1. I am more critical of myself than I used to be.
- 2. I criticize myself for all of my faults.
- 3. I blame myself for everything bad that happens.

9. Suicidal Thoughts or Wishes

- 0. I don't have any thoughts of killing myself.
- 1. I have thoughts of killing myself, but I would not carry them out.
- 2. I would like to kill myself.
- 3. I would kill myself if I had the chance.

10. Crying

- 0. I don't cry anymore than I used to.
- 1. I cry more than I used to.
- 2. I cry over every little thing.
- 3. I feel like crying, but I can't.

11. Agitation

- 0. I am no more restless or wound up than usual.
- 1. I feel more restless or wound up than usual.
- 2. I am so restless or agitated, it's hard to stay still.
- 3. I am so restless or agitated that I have to keep moving or doing something.

12. Loss of Interest

- 0. I have not lost interest in other people or activities.
- 1. I am less interested in other people or things than before.
- 2. I have lost most of my interest in other people or things.
- 3. It's hard to get interested in anything.

13. Indecisiveness

- 0. I make decisions about as well as ever.
- 1. I find it more difficult to make decisions than usual.
- 2. I have much greater difficulty in making decisions than I used to.
- 3. I have trouble making any decisions.

14. Worthlessness

- 0. I do not feel I am worthless.
- 1. I don't consider myself as worthwhile and useful as I used to.
- 2. I feel more worthless as compared to others.
- 3. I feel utterly worthless.

15. Loss of Energy

- 0. I have as much energy as ever.
- 1. I have less energy than I used to have.
- 2. I don't have enough energy to do very much.
- 3. I don't have enough energy to do anything.

16. Changes in Sleeping Pattern

- 0. I have not experienced any change in my sleeping.
- 1a I sleep somewhat more than usual.
- 1b I sleep somewhat less than usual.
- 2a I sleep a lot more than usual.
- 2b I sleep a lot less than usual.
- 3a I sleep most of the day.
- 3b I wake up 1-2 hours early and can't get back to sleep.

17. Irritability

- 0. I am not more irritable than usual.
- 1. I am more irritable than usual.
- 2. I am much more irritable than usual.
- 3. I am irritable all the time.

18. Changes in Appetite

- 0. I have not experienced any change in my appetite.
- 1a My appetite is somewhat less than usual.
- 1b My appetite is somewhat greater than usual.
- 2a My appetite is much less than before.
- 2b My appetite is much greater than usual.
- 3a I have no appetite at all.
- 3b I crave food all the time.

19. Concentration Difficulty

- 0. I can concentrate as well as ever.
- 1. I can't concentrate as well as usual.
- 2. It's hard to keep my mind on anything for very long.
- 3. I find I can't concentrate on anything.

20. Tiredness or Fatigue

- 0. I am no more tired or fatigued than usual.
- 1. I get more tired or fatigued more easily than usual.
- 2. I am too tired or fatigued to do a lot of the things I used to do.
- 3. I am too tired or fatigued to do most of the things I used to do.

21. Loss of Interest in Sex

- 0. I have not noticed any recent change in my interest in sex.
- 1. I am less interested in sex than I used to be.
- 2. I am much less interested in sex now.
- 3. I have lost interest in sex completely.

[Annexe B] : mail d'acceptation du paper + page de garde du paper.

- Mail d'acceptation pour la publication du paper :

----- Forwarded message -----

De : **CLEF 2020** <clef2020@easychair.org>

Date: jeu. 16 juil. 2020 à 15:09

Subject: erisk @ CLEF 2020 notification for paper 39

To: Fatima Boumahdi <f_boumahdi@esi.dz>

Dear Authors

We are glad to inform you that your working notes have been accepted to be published in the CLEF 2020 proceedings (CEUR workshop proceedings).

Find attached the reviews made by the referees. Please, take them into account when producing the final version of your paper.

You'll be required to upload your final version in easychair (more instructions about this will be sent by the proceedings editors).

Best Wishes

David, Fabio & Javier

SUBMISSION: 39

TITLE: USDB at eRisk 2020: Deep learning models to measure the Severity of the Signs of Depression using Reddit Posts

----- REVIEW 1 -----

SUBMISSION: 39

TITLE: USDB at eRisk 2020: Deep learning models to measure the Severity of the Signs of Depression using Reddit Posts

AUTHORS: Amina Madani, Fatima Boumahdi, Anfel Boukenaoui, Mohamed Chaouki Kritli and Hamza Hentabli

- Le paper :

USDB at eRisk 2020: Deep learning models to measure the Severity of the Signs of Depression using Reddit Posts

Amina MADANI¹, Fatima BOUMAHDI¹, Anfel BOUKENAOUTI¹, Mohamed Chaouki KRITLI¹, and Hamza HENTABLI²

¹ LRDSI Laboratory, BLIDA 1 University, Blida , Algeria
 {a_madani,f_boumahdi}@esi.dz, anfelboukenaoui@gmail.com,
 kritlichawki56@gmail.com

² Information Assurance and Security Research Group, Faculty of Computing,
 Universiti Teknologi Malaysia, Malaysia
 hentabli_hamza@yahoo.fr

Abstract. In this paper, we describe the participation of our USDB group (University of Saad Dahleb Blida) in the shared task T2 of the eRisk Lab at the CLEF 2020 workshop. This task focused on measuring the severity of the signs of depression from a thread of user posts. In response to this task, we study the performance of two different deep learning models (CNN and BiLSTM) in order to provide more perspectives for depression researches.

Keywords: Depression severity · Social networks · Natural language processing · Deep learning · CNN · BiLSTM · Sentiment analysis.

1 Introduction

Depression identification has been the subject of research of many fields, psychiatry, psychology, medicine and even sociolinguistics fields. Depression comes in different degrees and the examinations are usually done through one of the popular questionnaires used by psychologists, such as the Center of Epidemiologic Studies Depression Scale (CES-D) [26], Beck's Depression Inventory (BDI) [4] and Zung's Self-Rating Depression Scale (SDS) [38]. But, these examinations lack empirical data as they use the patient's observations or a third-party's ones which puts the results under the risk of flawed subjective human testing that can be manipulated easily, often with the purpose of gaining antidepressants or just to hide one's own depression from peers [21].

Twitter, Facebook and Reddit are different social media platforms that allow people to share their opinions and their personal thoughts. It has been proven

Copyright © 2020 for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0). CLEF 2020, 22-25 September 2020, Thessaloniki, Greece.

Références bibliographiques

1. organization, w.h. *Depression*. 2020; Available from: https://www.who.int/whr/2001/media_centre/press_release/fr/.
2. Nasukawa, T. and J. Yi, *Sentiment analysis: Capturing favorability using natural language processing*. 2003. 70-77.
3. Schouten, K. and F. Frasincar, *Survey on aspect-level sentiment analysis*. IEEE Transactions on Knowledge and Data Engineering, 2015. **28**(3): p. 813-830.
4. Cambria, E., *Affective computing and sentiment analysis*. IEEE intelligent systems, 2016. **31**(2): p. 102-107.
5. Yue, L., et al., *A survey of sentiment analysis in social media*. Knowledge and Information Systems, 2019: p. 1-47.
6. Katz, J.E. and R.E. Rice, *Public views of mobile medical devices and services: A US national survey of consumer sentiments towards RFID healthcare technology*. International journal of medical informatics, 2009. **78**(2): p. 104-114.
7. Miller, M., et al. *Sentiment Flow Through Hyperlink Networks*. in ICWSM. 2011.
8. Hong, Y. and S. Skiena. *The wisdom of bookies? sentiment analysis vs. the NFL point spread*. in *Proceedings of the international conference on Weblogs and Social media (icWSm-2010)*. 2010. Citeseer.
9. McGlohon, M., N. Glance, and Z. Reiter, *Star quality: Aggregating reviews to rank products and merchants*. 2010.
10. O'Connor, B., et al., *From tweets to polls: Linking text sentiment to public opinion time series*. Tepper School of Business, 2010: p. 559.
11. Rathi, R., *Effect of Cambridge Analytica's Facebook ads on the 2016 US Presidential election*. Towards Data Science, 2019.
12. McGlohon, M., N. Glance, and Z. Reiter, *Star quality: Aggregating reviews to rank products and merchants*. 2010.

Références bibliographiques

13. Liu, Y., et al. *ARSA: a sentiment-aware model for predicting sales performance using blogs*. in *Proceedings of the 30th annual international ACM SIGIR conference on Research and development in information retrieval*. 2007.
14. Petz, G., et al. *Opinion mining on the web 2.0—characteristics of user generated content and their impacts*. in *International Workshop on Human-Computer Interaction and Knowledge Discovery in Complex, Unstructured, Big Data*. 2013. Springer.
15. Joshi, M., et al. *Movie reviews and revenues: An experiment in text regression*. in *Human Language Technologies: The 2010 Annual Conference of the North American Chapter of the Association for Computational Linguistics*. 2010.
16. Asur, S. and B.A. Huberman. *Predicting the future with social media*. in *2010 IEEE/WIC/ACM international conference on web intelligence and intelligent agent technology*. 2010. IEEE.
17. Bollen, J., H. Mao, and X. Zeng, *Twitter mood predicts the stock market*. *Journal of computational science*, 2011. **2**(1): p. 1-8.
18. Troussas, C. and M. Virvou, *Advances in Social Networking-based Learning*.
19. Kawathekar, S.A. and D.M.M. Kshirsagar, *Movie Review analysis using Rule-Based & Support Vector Machines methods*. *IOSR Journal of Engineering*, 2012. **2**(3): p. 389-391.
20. Vohra, S. and J. Teraiya, *A comparative study of sentiment analysis techniques*. *Journal Jikrce*, 2013. **2**(2): p. 313-317.
21. Kalarikkal, S. and P. Remya. *Sentiment analysis and dataset collection: A comparative study*. in *2015 IEEE International Advance Computing Conference (IACC)*. 2015. IEEE.
22. Gao, S., J. Hao, and Y. Fu. *The application and comparison of web services for sentiment analysis in tourism*. in *2015 12th International Conference on Service Systems and Service Management (ICSSSM)*. 2015. IEEE.
23. Kaushik, C. and A. Mishra, *A scalable, lexicon based technique for sentiment analysis*. *arXiv preprint arXiv:1410.2265*, 2014.

Références bibliographiques

24. Devika, M., C. Sunitha, and A. Ganesh, *Sentiment analysis: a comparative study on different approaches*. Procedia Computer Science, 2016. **87**: p. 44-49.
25. Cambria, E., et al., *Knowledge-based approaches to concept-level sentiment analysis*. IEEE intelligent systems, 2013. **28**(2): p. 12-14.
26. Kanakaraj, M. and R.M.R. Guddeti. *NLP based sentiment analysis on Twitter data using ensemble classifiers*. in *2015 3rd International Conference on Signal Processing, Communication and Networking (ICSCN)*. 2015. IEEE.
27. Maynard, D. and A. Funk. *Automatic detection of political opinions in tweets*. in *Extended Semantic Web Conference*. 2011. Springer.
28. Li, G. and F. Liu, *Application of a clustering method on sentiment analysis*. Journal of Information Science, 2012. **38**(2): p. 127-139.
29. Moraes, R., J.F. Valiati, and W.P.G. Neto, *Document-level sentiment classification: An empirical comparison between SVM and ANN*. Expert Systems with Applications, 2013. **40**(2): p. 621-633.
30. Zhang, C., et al., *Sentiment analysis of Chinese documents: From sentence to document level*. Journal of the American Society for Information Science and Technology, 2009. **60**(12): p. 2474-2487.
31. Wiebe, J., R. Bruce, and T.P. O'Hara. *Development and use of a gold-standard data set for subjectivity classifications*. in *Proceedings of the 37th annual meeting of the Association for Computational Linguistics*. 1999.
32. Liu, B., *Sentiment analysis and opinion mining*. Synthesis lectures on human language technologies, 2012. **5**(1): p. 1-167.
33. Bar-Haim, R., et al. *Identifying and following expert investors in stock microblogs*. in *Proceedings of the 2011 Conference on Empirical Methods in Natural Language Processing*. 2011.
34. Medhat, W., A. Hassan, and H. Korashy, *Sentiment analysis algorithms and applications: A survey*. Ain Shams engineering journal, 2014. **5**(4): p. 1093-1113.

Références bibliographiques

35. Rajkomar, A., J. Dean, and I. Kohane, *Machine learning in medicine*. New England Journal of Medicine, 2019. **380**(14): p. 1347-1358.
36. Pang, B., L. Lee, and S. Vaithyanathan, *Thumbs up? Sentiment classification using machine learning techniques*. arXiv preprint cs/0205070, 2002.
37. Otter, D.W., J.R. Medina, and J.K. Kalita, *A survey of the usages of deep learning for natural language processing*. IEEE Transactions on Neural Networks and Learning Systems, 2020.
38. Qiu, X., et al., *Pre-trained models for natural language processing: A survey*. arXiv preprint arXiv:2003.08271, 2020.
39. Baroni, M., *Linguistic generalization and compositionality in modern artificial neural networks*. Philosophical Transactions of the Royal Society B, 2020. **375**(1791): p. 20190307.
40. Young, T., et al., *Recent trends in deep learning based natural language processing*. iee Computational intelligence magazine, 2018. **13**(3): p. 55-75.
41. Henderson, M., et al., *Efficient natural language response suggestion for smart reply*. arXiv preprint arXiv:1705.00652, 2017.
42. Oord, A.v.d., et al., *Wavenet: A generative model for raw audio*. arXiv preprint arXiv:1609.03499, 2016.
43. Wojna, Z., et al. *Attention-based extraction of structured information from street view imagery*. in *2017 14th IAPR International Conference on Document Analysis and Recognition (ICDAR)*. 2017. IEEE.
44. Khrabrov, K., et al., *The cornucopia of meaningful leads: Applying deep adversarial autoencoders for new molecule development in oncology*. 2017.
45. Polson, N.G. and V. Sokolov, *Deep learning: A Bayesian perspective*. Bayesian Analysis, 2017. **12**(4): p. 1275-1304.
46. Dixon, M.F., N.G. Polson, and V.O. Sokolov, *Deep learning for spatio-temporal modeling: Dynamic traffic flows and high frequency trading*. arXiv preprint arXiv:1705.09851, 2017.

Références bibliographiques

47. Heaton, J.B., N.G. Polson, and J.H. Witte, *Deep learning for finance: deep portfolios*. Applied Stochastic Models in Business and Industry, 2017. **33**(1): p. 3-12.
48. Sokolov, V., *Discussion of 'Deep learning for finance: deep portfolios'*. Applied Stochastic Models in Business and Industry, 2017. **33**(1): p. 16-18.
49. Feng, G., N. Polson, and J. Xu, *Deep learning in characteristics-sorted factor models*. Available at SSRN 3243683, 2019.
50. Feng, G., J. He, and N.G. Polson, *Deep learning for predicting asset returns*. arXiv preprint arXiv:1804.09314, 2018.
51. Kamath, U., J. Liu, and J. Whitaker, *Deep learning for nlp and speech recognition*. Vol. 84. 2019: Springer.
52. Abdel-Hamid, O., L. Deng, and D. Yu. *Exploring convolutional neural network structures and optimization techniques for speech recognition*. in *Interspeech*. 2013.
53. Chouhan, N. and A. Khan, *Network anomaly detection using channel boosted and residual learning based deep convolutional neural network*. Applied Soft Computing, 2019. **83**: p. 105612.
54. Gu, J., et al., *Recent advances in convolutional neural networks*. Pattern Recognition, 2018. **77**: p. 354-377.
55. Mittal, V., D. Gangodkar, and B. Pant. *Exploring The Dimension of DNN Techniques For Text Categorization Using NLP*. in *2020 6th International Conference on Advanced Computing and Communication Systems (ICACCS)*. 2020. IEEE.
56. Kamath, U., J. Liu, and J. Whitaker, *Convolutional Neural Networks*, in *Deep Learning for NLP and Speech Recognition*. 2019, Springer. p. 263-314.
57. Ahonen, L., *Violence and Mental Illness: An Overview*. 2019: Springer.
58. Errecalde, M., et al. *Temporal Variation of Terms as Concept Space for Early Risk Prediction*. in *CLEF*. 2017.
59. Funez, D.G., et al. *UNSL's participation at eRisk 2018 Lab*. in *CLEF (Working Notes)*. 2018.

Références bibliographiques

60. Arlington, V., *American Psychiatric Publishing*. Diagnostic and statistical manual of mental disorders (5-th ed), 2013: p. 5-25.
61. Hewitt, M. and J.H. Rowland, *Mental health service use among adult cancer survivors: analyses of the National Health Interview Survey*. Journal of Clinical Oncology, 2002. **20**(23): p. 4581-4590.
62. Whittaker, S., E. Isaacs, and V. O'Day, *Widening the net: workshop report on the theory and practice of physical and network communities*. ACM SIGCHI Bulletin, 1997. **29**(3): p. 27-30.
63. Caplan, S.E. and J.S. Turner, *Bringing theory to research on computer-mediated comforting communication*. Computers in human behavior, 2007. **23**(2): p. 985-998.
64. De Choudhury, M. and S. De. *Mental health discourse on reddit: Self-disclosure, social support, and anonymity*. in *Eighth international AAAI conference on weblogs and social media*. 2014.
65. De Choudhury, M. and E. Kıcıman. *The language of social support in social media and its effect on suicidal ideation risk*. in *Proceedings of the... International AAAI Conference on Weblogs and Social Media. International AAAI Conference on Weblogs and Social Media*. 2017. NIH Public Access.
66. De Choudhury, M. and E. Kıcıman. *The language of social support in social media and its effect on suicidal ideation risk*. in *Proceedings of the... International AAAI Conference on Weblogs and Social Media. International AAAI Conference on Weblogs and Social Media*. 2017. NIH Public Access.
67. Becker, B. and G. Mark, *Social conventions in computermediated communication: A comparison of three online shared virtual environments*, in *The social life of avatars*. 2002, Springer. p. 19-39.
68. Eysenbach, G., et al., *Health related virtual communities and electronic support groups: systematic review of the effects of online peer to peer interactions*. Bmj, 2004. **328**(7449): p. 1166.

Références bibliographiques

69. Postmes, T., R. Spears, and M. Lea, *The formation of group norms in computer-mediated communication*. Human communication research, 2000. **26**(3): p. 341-371.
70. Ferrara, K., *Accommodation in therapy*. Contexts of accommodation: Developments in applied sociolinguistics, 1991: p. 187-222.
71. Beaudoin, C.E. and C.-C. Tao, *Benefiting from social capital in online support groups: An empirical study of cancer patients*. CyberPsychology & Behavior, 2007. **10**(4): p. 587-590.
72. O'Keeffe, G.S. and K. Clarke-Pearson, *The impact of social media on children, adolescents, and families*. Pediatrics, 2011. **127**(4): p. 800-804.
73. De Choudhury, M., et al., *Predicting depression via social media*. Icwsm, 2013. **13**: p. 1-10.
74. O'dea, B., et al., *Detecting suicidality on Twitter*. Internet Interventions, 2015. **2**(2): p. 183-188.
75. Zhang, L., et al. *Using linguistic features to estimate suicide probability of Chinese microblog users*. in *International Conference on Human Centered Computing*. 2014. Springer.
76. Nguyen, T., et al., *Affective and content analysis of online depression communities*. IEEE Transactions on Affective Computing, 2014. **5**(3): p. 217-226.
77. Kaslow, N.J., et al., *An empirical study of the psychodynamics of suicide*. Journal of the American Psychoanalytic Association, 1998. **46**(3): p. 777-796.
78. Levine, R.S., et al., *Firearms, youth homicide, and public health*. Journal of health care for the poor and underserved, 2012. **23**(1): p. 7.
79. Hall, R.C., D.E. Platt, and R.C. Hall, *Suicide risk assessment: a review of risk factors for suicide in 100 patients who made severe suicide attempts: evaluation of suicide risk in a time of managed care*. Psychosomatics, 1999. **40**(1): p. 18-27.
80. Mann, J.J., et al., *Suicide prevention strategies: a systematic review*. Jama, 2005. **294**(16): p. 2064-2074.

Références bibliographiques

81. Akram, W., *A Study on Positive and Negative Effects of Social Media on Society*. International Journal of Computer Sciences and Engineering, 2018. **5**.
82. Shickel, B., et al. *Self-reflective sentiment analysis*. in *Proceedings of the Third Workshop on Computational Linguistics and Clinical Psychology*. 2016.
83. Singh, D. and A. Wang, *Detecting Depression Through Tweets*. Stanford University CA 9430: p. 1-9.
84. Wang, Y.-T., H.-H. Huang, and H.-H. Chen. *A Neural Network Approach to Early Risk Detection of Depression and Anorexia on Social Media Text*. in *CLEF (Working Notes)*. 2018.
85. Losada, D.E., F. Crestani, and J. Parapar. *Early detection of risks on the Internet: an exploratory campaign*. in *European Conference on Information Retrieval*. 2019. Springer.
86. Beck, A.T., R.A. Steer, and M.G. Carbin, *Psychometric properties of the Beck Depression Inventory: Twenty-five years of evaluation*. Clinical psychology review, 1988. **8**(1): p. 77-100.
87. Wilson, T., J. Wiebe, and P. Hoffmann. *Recognizing contextual polarity in phrase-level sentiment analysis*. in *Proceedings of human language technology conference and conference on empirical methods in natural language processing*. 2005.
88. Van Rijen, P., et al. *A Data-Driven Approach for Measuring the Severity of the Signs of Depression using Reddit Posts*. in *CLEF (Working Notes)*. 2019.
89. Kraskov, A., H. Stögbauer, and P. Grassberger, *Erratum: estimating mutual information [Phys. Rev. E 69, 066138 (2004)]*. Physical Review E, 2011. **83**(1): p. 019903.
90. Pennington, J., R. Socher, and C.D. Manning. *Glove: Global vectors for word representation*. in *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*. 2014.
91. Abed-Esfahani, P., et al. *Transfer Learning for Depression: Early Detection and Severity Prediction from Social Media Postings*. in *CLEF (Working Notes)*. 2019.
92. Fellbaum, C., *WordNet*. The encyclopedia of applied linguistics, 2012.

Références bibliographiques

93. Eichstaedt, J.C., et al., *Facebook language predicts depression in medical records*. Proceedings of the National Academy of Sciences, 2018. **115**(44): p. 11203-11208.
94. Burdisso, S., M.L. Errecalde, and M. Montes y Gómez. *Towards Measuring the Severity of Depression in Social Media via Text Classification*. in *XXV Congreso Argentino de Ciencias de la Computación (CACIC)*(Universidad Nacional de Río Cuarto, Córdoba, 14 al 18 de octubre de 2019). 2019.
95. Burdisso, S.G., M. Errecalde, and M. Montes-y-Gómez, *A text classification framework for simple and effective early depression detection over social media streams*. Expert Systems with Applications, 2019. **133**: p. 182-197.
96. Losada, D.E., F. Crestani, and J. Parapar. *Overview of eRisk: early risk prediction on the internet*. in *International Conference of the Cross-Language Evaluation Forum for European Languages*. 2018. Springer.
97. Abed-Esfahani, P., et al. *Transfer Learning for Depression: Early Detection and Severity Prediction from Social Media Postings*. in *CLEF (Working Notes)*. 2019.
98. Trifan, A. and J.L. Oliveira. *BioInfo@ UAVR at eRisk 2019: delving into Social Media Texts for the Early Detection of Mental and Food Disorders*. in *CLEF (Working Notes)*. 2019.
99. Mikolov, T., et al. *Distributed representations of words and phrases and their compositionality*. in *Advances in neural information processing systems*. 2013.
100. Rawat, W. and Z. Wang, *Deep convolutional neural networks for image classification: A comprehensive review*. Neural computation, 2017. **29**(9): p. 2352-2449.
101. Beck, A.T., R.A. Steer, and M.G. Carbin, *Psychometric properties of the Beck Depression Inventory: Twenty-five years of evaluation*. Clinical psychology review, 1988. **8**(1): p. 77-100.