

Modèles de classification et de prédiction de performance des
signaux optiques WDM basés sur les algorithmes
d'apprentissage machine

par

Stéphanie ALLOGBA

THÈSE PRÉSENTÉE À L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
COMME EXIGENCE PARTIELLE À L'OBTENTION DU
DOCTORAT EN GÉNIE
Ph. D.

MONTREAL, LE 17 NOVEMBRE 2021

ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
UNIVERSITÉ DU QUÉBEC

©Tous droits réservés, Stéphanie Allogba, 2021

©Tous droits réservés

Cette licence signifie qu'il est interdit de reproduire, d'enregistrer ou de diffuser en tout ou en partie, le présent document. Le lecteur qui désire imprimer ou conserver sur un autre média une partie importante de ce document, doit obligatoirement en demander l'autorisation à l'auteur.

PRÉSENTATION DU JURY

CETTE THÈSE A ÉTÉ ÉVALUÉE

PAR UN JURY COMPOSÉ DE :

Mme Christine Tremblay, directrice de thèse
Département de génie électrique à l'École de technologie supérieure

M. Robert Sabourin, président du jury
Département de génie des systèmes à l'École de technologie supérieure

M. Zbigniew Dziong, membre du jury
Département de génie électrique à l'École de technologie supérieure

Mme Halima Elbiaze, examinatrice externe
Département d'informatique à l'Université du Québec à Montréal

ELLE A FAIT L'OBJET D'UNE SOUTENANCE DEVANT JURY ET PUBLIC

LE 21 OCTOBRE 2021

À L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE

REMERCIEMENTS

Avant tout, je profite de ces quelques lignes pour adresser mes remerciements à tous ceux qui m'ont aidé à la réalisation de ce projet.

Ainsi, mes remerciements s'adressent à ma directrice de recherche, Professeure Christine Tremblay pour m'avoir offert la possibilité de réaliser mon projet doctoral au sein de son équipe de recherche. J'ai ainsi pu avoir accès à un domaine porteur dans le monde de la télécommunication et de l'apprentissage machine. De plus, son encadrement, sa rigueur et sa disponibilité et ses conseils avisés tout au long de mon projet doctoral m'ont permis de finaliser ce rapport et m'outilleront pendant tout le reste de mon parcours professionnel.

Mes remerciements vont également à l'endroit des membres du jury pour leur commentaire sur mon travail.

Je tiens à remercier également toute l'équipe du Laboratoire de technologies de réseaux avec qui j'ai eu le plaisir de travailler. Un merci spécial à Laure et Sandra pour avoir été de bonnes conseillères et des soutiens constants.

Merci également à tous mes amis et ma famille de cœur, en particulier à Landry, Adnette et Aguessy, de m'avoir soutenu pendant toute la durée de mon projet doctoral.

À ma mère, Valérie et mon très cher clan Allogba, je ne saurai vous remercier pour votre soutien depuis le début de mes études. Merci d'avoir sans cesse été présents pour moi, à chaque étape que j'ai eu à traverser jusqu'à présent.

Je ne saurai terminer sans exprimer ma reconnaissance envers mon époux Parfait. Merci d'avoir cru en moi et de m'avoir soutenu dans les moments où je pensais lâcher. Merci pour ta patience et ton encouragement à chaque étape que je franchissais. Ahossou Tché, tu as été pour

VI

moi une force. Ta patience, ta détermination et tes encouragements m'ont donné l'envie d'aller le plus loin possible et de donner le meilleur de moi pour finaliser mon doctorat.

À vous qui n'avez sans cesse veillé sur moi, mes papas Alain et Dada Gnankan, je ne saurai vous remercier pour toutes les valeurs que vous m'avez inculquées dès le bas âge. J'espère vous rendre continuellement fier de mon parcours.

Modèles de classification et de prédiction de performance des signaux optiques WDM basés sur les algorithmes d'apprentissage machine

Stéphanie ALLOGBA

RÉSUMÉ

La croissance constante du trafic contraint les opérateurs de télécommunications à déployer des systèmes de transmission optique WDM (*Wavelength-Division Multiplexing*) ayant des débits, une capacité et une flexibilité en constante évolution. Cependant, ces systèmes transportant une multitude d'applications (vidéo, *cloud*, Internet des objets ...), l'impact des dégradations de leur performance devient de plus en plus important. Une solution potentielle serait d'implémenter des outils de contrôle et de gestion de réseau proactifs pouvant exploiter les métriques de performance collectées sur le réseau. Dans ce contexte, l'apprentissage machine se positionne progressivement comme une solution prometteuse pour gérer la complexité et l'évolutivité des réseaux optiques hétérogènes. En effet, les outils basés sur les techniques d'apprentissage machine seraient très utiles dans la couche physique de réseaux optiques dans un contexte de réseau dynamique défini par logiciel (*Software Defined Network*, SDN), en particulier pour la prédiction des performances.

Cette thèse vise donc à proposer de nouvelles méthodes basées sur des algorithmes d'apprentissage machine pouvant être intégrés dans les outils de gestion des réseaux optiques afin de contrôler les performances des connexions optiques qui transportent un trafic pouvant atteindre plusieurs Térabits par seconde. En outre, chacune des méthodes proposées est validée par des simulations utilisant des données de performance réelles recueillies dans des liaisons optiques de réseaux de production. Visant à répondre à des problèmes différents, ces nouvelles méthodes proposées sont présentées dans trois parties distinctes de la thèse.

Dans la première partie, la première méthode développée est un outil de classification construit à partir des données de BER (*Bit Error Rate*) d'une connexion optique (*lightpath*) déployée dans un réseau de production pendant une période de 31 jours. Cet outil utilise l'algorithme des plus proches voisins (*K-nearest neighbor*, KNN). Il a pour but de caractériser le niveau de qualité du BER et répond donc au problème d'estimation de la qualité de transmission d'une connexion optique. De plus, dans cette partie, un module permettant d'analyser l'impact des attributs utilisés pour estimer la qualité de transmission d'une connexion optique a été implémenté. Ce module a été construit à partir de l'algorithme à support vecteurs machine (*Support Vector Machine*, SVM).

Dans la seconde partie, les modèles proposés ont pour but de répondre au problème de prédiction des métriques de performance d'une connexion optique, notamment le BER et le rapport signal-sur-bruit (SNR). Ces modèles sont construits en utilisant deux variantes des réseaux de neurones (*Long Short-Term Memory*, LSTM et *Gated Recurrent Unit*, GRU). Cette partie est divisée en trois sections. La première section présente des modèles univariés de

prédiction du SNR d'une connexion optique utilisant les données historiques de BER provenant d'un réseau de production. La deuxième section présente les modèles multivariés de prédiction de SNR d'une connexion optique utilisant à la fois les données historiques de BER ainsi que d'autres attributs comme la puissance optique du canal WDM, la température, etc. La troisième section vise à explorer le potentiel de l'apprentissage par transfert en évaluant les modèles multivariés et univariés de prédiction de SNR sur des ensembles de données de terrain provenant de connexions optiques différentes.

Finalement, dans la troisième partie, un modèle de détection d'anomalies de performance, basé sur un algorithme SVM et l'approche de l'écart interquartile (*Inter Quartile Range*, IQR), est proposé. Dans un premier temps, le modèle permet de définir et d'extraire les attributs permettant de caractériser les anomalies observées dans les métriques de performance de connexions optiques déployées dans un réseau de production. Dans un deuxième temps, le modèle est utilisé pour la détection anticipée des anomalies dans les séries temporelles de SNR recueillies dans le réseau de production.

Mots-clés : apprentissage machine, apprentissage par transfert, détection anticipée d'anomalies, modèle univarié, modèle multivarié, prédiction des performances, qualité de transmission, réseaux de neurones

Classification and prediction models based on machine learning algorithms for lightpath quality of transmission (QoT) estimation and prediction

Stéphanie ALLOGBA

ABSTRACT

The constant growth of the traffic is forcing telecom operators to deploy Wavelength-Division Multiplexing (WDM) optical transmission systems with constantly evolving data rates, capacity, and flexibility. However, as these systems carry a multitude of applications (video, cloud, Internet of Things ...); the impact of performance degradations is becoming increasingly important. A potential solution would be to implement proactive network monitoring and management tools that can exploit performance metrics collected on the network. In this context, Machine Learning (ML) is gradually positioning itself as a promising solution to manage the complexity and scalability of heterogeneous optical networks. Indeed, tools based on ML techniques would be very useful in the physical layer of optical networks in a Software Defined Network (SDN) context, in particular for performance prediction.

This thesis aims at proposing new methods based on ML algorithms that can be integrated into network control system in order to manage the performance of the lightpaths deployed in optical networks which carry traffic up to several Terabits per second. In addition, each of the proposed methods is validated by simulations using field performance data collected from lightpaths deployed in production networks. Aiming at addressing different problems, these new proposed methods are presented in three parts of the thesis.

In the first part, the proposed ML method is a classifier built with Bit Error Rate (BER) data from a lightpath carried in a production network during a 31-day observation period. The classifier, based on the K-nearest neighbor (K-NN) algorithm, aims at characterizing the BER quality level and addresses the problem of estimating the quality of transmission (QoT) of an established lightpath. Moreover, in this part, a module for assessing the impact of different features used for lightpath QoT estimation has been implemented using the Support Vector Machine (SVM) algorithm.

In the second part, the proposed ML models aim at answering the problem of predicting the performance, namely the BER and signal-to-noise ratio (SNR) of established lightpaths, based on historical field data. The models are built using two variants of neural networks, namely Long Short-Term Memory (LSTM) and Gated Recurrent Unit (GRU). This part is divided into three sections. The first section presents univariate SNR forecast models using historical lightpath BER data only to implement the SNR forecast module. The second section presents multivariate SNR forecast models using both historical BER data and additional features such as WDM channel power, temperature and others in the forecast module. The third section explores the potential of transfer learning by evaluating the multivariate and univariate SNR forecast models with field data from a different set of deployed lightpaths.

Finally, in the third part, the proposed ML method is an anomaly detection tool based on the SVM and the Inter Quartile Range (IQR) methods. First, the tool defines and extracts the features that characterize the anomalies observed in the SNR time series of several lightpaths carried in a production network. Then, the tool performs an early detection of the anomalies in the SNR time series. The anomaly detection tool is tested on field performance data collected in a production network.

Keywords: early anomaly detection, machine learning, multivariate model, neural networks, performance prediction, quality of transmission, transfer learning, univariate model.

TABLE DES MATIÈRES

	Page
INTRODUCTION 1	
0.1	Problématique1
0.2	Contributions.....2
0.2.1	Estimateur de performance d'une connexion optique 3
0.2.2	Prédicteur de performance d'une connexion optique 3
0.2.3	Détecteur des anomalies de performance d'une connexion optique..... 3
0.2.4	Contributions scientifiques 4
0.3	Organisation.....5
CHAPITRE 1 REVUE DE LA LITTÉRATURE.....7	
1.1	Introduction.....7
1.2	Réseaux optiques : présentation générale8
1.2.1	Description générale 8
1.2.2	Effets de dégradation du signal optique..... 10
1.2.2.1	Atténuation..... 11
1.2.2.2	Dispersion chromatique 11
1.2.2.3	Dispersion des modes de polarisation..... 12
1.2.2.4	Effets non linéaires 13
1.2.3	Paramètres de performance..... 14
1.2.3.1	Taux d'erreur binaire 14
1.2.3.2	Rapport signal-sur-bruit optique 14
1.3	Apprentissage machine16
1.3.1	Présentation générale 16
1.3.2	Méthodologie générale..... 18
1.3.2.1	Description des termes..... 18
1.3.2.2	Présentation des méthodes d'apprentissage machine 20
1.4	Apprentissage machine dans les télécommunications optiques.....22
1.4.1	Estimation de la qualité de transmission..... 23
1.4.2	Détection des pannes..... 27
CHAPITRE 2 CLASSIFICATEUR DE PERFORMANCE D'UNE CONNEXION OPTIQUE.....31	
2.1	Introduction.....31
2.2	Classificateur de performance utilisant l'algorithme K-NN.....32
2.2.1	Prétraitement des données..... 33
2.2.2	Implémentation de l'algorithme de classification..... 39
2.3	Étude de l'impact des attributs dans les classificateurs44
2.3.1	Présentation de l'estimateur de qualité 44
2.3.2	Étude de l'impact des attributs..... 46

2.4	Conclusion	50
CHAPITRE 3 Outils de Prédiction des Performances d'une Connexion Optique.....53		
3.1	Introduction.....	53
3.2	Objectif	54
3.3	Génération de la base de données	56
3.3.1	Caractéristiques des deux connexions optiques du réseau de CANARIE	56
3.3.2	Prétraitement des données.....	59
3.4	Modèles de prédiction de performance d'une connexion optique.....	63
3.4.1	Description des modèles de prédiction de SNR.....	63
3.4.2	Implémentation des modèles de prédiction du SNR.....	66
3.5	Évaluation de performance	70
3.5.1	Modèle LSTM-1A : résultats	71
3.5.2	Modèle GRU-U-1B : résultats	73
3.5.3	Modèles LSTM-U-2A et GRU-U-2A : résultats	75
3.6	Conclusion	79
CHAPITRE 4 Modèles Multivariés pour la Prédiction de Performance d'une Connexion Optique.....81		
4.1	Introduction.....	81
4.2	Objectif	82
4.3	Collecte et prétraitement des données.....	85
4.3.1	Présentation des bases de données.....	85
4.3.2	Analyse statistique	88
4.4	Modèles de prédiction multivariés.....	91
4.4.1	Présentation des modèles multivariés et optimisation des hyper-paramètres	91
4.4.2	Évaluation des performances	93
4.5	Modèle de prédiction multivarié avec réduction des attributs.....	97
4.5.1	Ingénierie des attributs et présentations des modèles réduits	97
4.5.2	Évaluation des performances des modèles réduits.....	98
4.6	Généralisation des modèles de prédiction utilisant l'apprentissage par transfert	101
4.6.1	Description de l'approche d'apprentissage par transfert	102
4.6.2	Étapes d'implémentation	104
4.6.3	Évaluation des performances	105
4.7	Conclusion	109
CHAPITRE 5 Modèle de Détection Anticipée des Anomalies de Performance de Connexions Optiques.....113		
5.1	Introduction.....	113
5.2	Définition du problème	115
5.2.1	Présentation des travaux	115

5.2.2	Objectif	117
5.3	Module de définition des anomalies	120
5.3.1	Description du réseau.....	120
5.3.2	Définition des anomalies.....	121
5.3.3	Identification des attributs.....	124
5.4	Module de détection des anomalies	125
5.4.1	Ingénierie des attributs.....	125
5.4.1.1	Phase de nettoyage des attributs	126
5.4.1.2	Sélection des attributs	127
5.4.2	Modèles de détection des anomalies.....	132
5.4.2.1	Implémentation de l'algorithme SVM	132
5.4.2.2	Évaluation des performances	135
5.5	Conclusion	140
CONCLUSION		143
RECOMMANDATIONS		147
LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES.....		149

LISTE DES TABLEAUX

	Page
Tableau 2.1	Évaluation des performances du classificateur de BER42
Tableau 2.2	Paramètre de l’outil de générateur de données46
Tableau 2.3	Évaluation de performance de l’estimateur de qualité.....49
Tableau 3.1	Caractéristiques des deux bases de données de terrain de deux connexions optiques utilisées pour les modèles de prédiction de performance56
Tableau 3.2	Résultats des tests de stationnarité.....62
Tableau 3.3	Hyper-paramètres des modèles de prédiction de SNR68
Tableau 4.1	Caractéristiques des trois bases de données de terrain issues de trois connexions optiques utilisées pour les modèles de prédiction de performance86
Tableau 4.2	Résumé statistique des données de SNR, P_{RX} et DGD pour chaque connexion.....88
Tableau 4.3	Présentation des hyper-paramètres des modèles.....92
Tableau 4.4	Analyse de corrélation des attributs de KB-197
Tableau 4.5	Analyse comparative des modèles univariés et multivariés100
Tableau 5.1	Présentation des meilleurs attributs sélectionnés à partir des méthodes de sélection d’attributs131
Tableau 5.2	Ensemble des hyper-paramètres optimaux134
Tableau 5.3	Présentation des performances des modèles de détection des anomalies136
Tableau 5.4	Synthèse des résultats138

LISTE DES FIGURES

	Page
Figure 1.1	Schéma général d'un réseau optique.....10
Figure 1.2	Étapes d'identification de la méthode d'apprentissage machine16
Figure 2.1	Étapes d'implémentation du classificateur de performance33
Figure 2.2	Connexion optique du réseau CANARIE34
Figure 2.3	Évolution temporelle du BER ; Probabilité de distribution du BER35
Figure 2.4	Exemple de représentation du seuil pour l'identification des étiquettes....36
Figure 2.5	Évolution de la température du pré-FEC BER en fonction du temps durant une période d'observation de 22 jours.....37
Figure 2.6	Probabilité de distribution du BER en fonction des périodes de la journée38
Figure 2.7	Principe de fonctionnement de l'algorithme K-NN.....40
Figure 2.8	Variation du taux de précision en fonction du nombre de voisins K.....41
Figure 2.9	Matrice de confusion du classificateur de BER utilisant la méthode par validation croisée à 10 plis.....43
Figure 2.10	Taux de précision en fonction du choix des attributs48
Figure 3.1	Procédure d'implémentation des modèles prédictifs55
Figure 3.2	Évolution temporelle du SNR de la connexion optique A (base de données A) durant la période d'observation de 13 mois58
Figure 3.3	Évolution temporelle du SNR de la connexion optique B (base de données B) durant la période d'observation de 18 mois.....58
Figure 3.4	Topologie des réseaux LSTM et GRU.....64
Figure 3.5	Présentation du RMSE en fonction de l'horizon pour la méthode LSTM-U-1A.....71

Figure 3.6	SNR prédit vs SNR observé (connexion A, LSTM-U-1A) pour les périodes de : (a) 5 janvier ; (b) 17 janvier.....	72
Figure 3.7	RMSE en fonction de l'horizon de prédiction pour les modèles GRU-U-1B, encodeur-décodeur LSTM et naïf	73
Figure 3.8	SNR prédit vs SNR observé (connexion B, GRU-U-1B) pour la période d'observation du 29 avril 2018 au 30 avril 2018	74
Figure 3.9	Performance des modèles LSTM-U-2A et GRU-U-2A : (a) RMSE en fonction de l'horizon ; (b) R2 en fonction de l'horizon.....	76
Figure 3.10	SNR prédit vs SNR observé (connexion A, LSTM-U-2A et GRU-U-2A) : (a) horizon de 24 heures ; (b) horizon de 48 heures ; (c) horizon de 96 heures	78
Figure 4.1	Topologie de la portion du réseau à l'étude (distances approximatives)...	85
Figure 4.2	Probabilité de distribution du SNR pour : (a) lightpath-1 ; (b) lightpath-2 ; (c) lightpath-3	87
Figure 4.3	Performance des modèles multivariés LSMT-M-1 et GRU-M-1 utilisant l'ensemble de test de KB-1 : (a) RMSE ; (b) AME ; (c) score R^2 ; (d) Temps de calcul.....	94
Figure 4.4	SNR prédit vs SNR observé (connexion 1, LSTM-M-1) : (a) horizon de 1 heure ; (b) horizon de 24 heures ; (c) horizon de 96 heures.....	96
Figure 4.5	Performance des modèles multivariés utilisant l'ensemble de test de KB-1 réduit : (a) RMSE ; (b) AME ; (c) score R^2	99
Figure 4.6	Présentation du concept d'apprentissage par transfert.....	103
Figure 4.7	Performance des modèles LSTM-M-2 et GRU-U-2A pour la base de données KB-3 (connexion optique de direction opposée sur la même route) : (a) RMSE ; (b) AME – Performances des modèles pour la base de données KB-2 (connexion optique utilisant la même fibre optique sur une portion de la même route) : (c) RMSE ; (d) AME	106
Figure 4.8	SNR prédit vs SNR observé pour l'horizon de 1 heure, 4 heures et 96 heures utilisant le modèle multivarié LSTM basé sur l'apprentissage par transfert : (a) connexion optique cible lightpath-3 (base de données KB-3) ; (b) connexion optique cible lightpath-2 (base de données KB-2)	108

Figure 5.1	Étapes d'implémentation du modèle de détection des anomalies dans les données de SNR de connexions optiques.....	119
Figure 5.2	Présentation de l'approche IQR pour la détermination des anomalies ...	122
Figure 5.3	(a) Identification des anomalies des 8 connexions optiques en utilisant l'approche IQR ; (b) Exemple d'anomalies observées sur les données de SNR d'une des 8 connexions optiques.....	123
Figure 5.4	Résultats de l'analyse de corrélation des attributs pour le modèle de détection d'anomalies	126
Figure 5.5	Processus de la méthode Boruta	128
Figure 5.6	Processus de la méthode RF	129
Figure 5.7	Processus du modèle RFE.....	130
Figure 5.8	Fonctionnement général de l'algorithme SVM.....	133
Figure 5.9	Matrices de confusion pour les modèles : (a) SVM_All ; (b) SVM_Boruta ; (c) SVM_RF ; (d) SVM_RFE	135

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

ACP	Analyse en composantes principales
AME	Absolute Maximum Error
ANN	Artificial Neural Network
ARIMA	AutoRegressive Moving Average
ASE	Amplified Spontaneous Emission
BER	Bit Error Rate
BPSK	Binary Phase Shift Keying
DEMUX	Demultiplexer
DGD	Differential Group Delay
DP	Dual Polarization
DSP	Digital Signal Processing
EDFA	Erbium Doped Fibre Amplifier
EM	Expectation Maximization
FEC	Forward Error Correction
FWM	Four Wave Mixing
GRU	Gated Recurrent Unit
IP	Internet Protocol
IQR	Interquartile Range
K-NN	K-Nearest Neighbor
LSTM	Long Short-Term Memory
LOESS	Locally Estimated Scatterplot Smoothing

MAE	Mean Absolute Error
MUX	Multiplexer
NARNN	Nonlinear AutoRegressive Neural Network
NN	Neural Network
OSNR	Optical Signal Noise Ratio
PCA	Principal Component Analysis
PDL	Polarization Dispersion Loss
PM	Phase Modulation
PMD	Polarization Mode Dispersion
QAM	Quadrature Amplitude Modulation
QPSK	Quadrature Phase-Shift Keying
RAM	Random access memory
RF	Random Forest
RFE	Recursive Feature Elimination
RMSE	Root Mean Square Error
ROADM	Reconfigurable Optical Add-Drop Multiplexer
ROC	Receiver Operating Characteristic
SBS	Stimulated Brillouin Scattering
SDN	Software Defined Network
SNR	Signal-to-Noise Ratio
SPM	Self-Phase Modulation
SRS	Stimulated Raman Scattering

STL	Seasonal and Trend decomposition using LOESS
SVM	Support Vector Machine
Tanh	Tangente hyperbolique
WDM	Wavelength-Division Multiplexing
WSS	Wavelength Selective Switch
XPM	Cross-phase Modulation

LISTE DES SYMBOLES ET UNITÉS DE MESURE

FRÉQUENCE

GHz Gigahertz

THz Téraherzt

UNITÉS DE TEMPS

h heure

min minute

s seconde

ps picoseconde

C Celsius

UNITÉS DE DÉBIT

Gbit/s Gigabit par seconde

Mbit/s Mégabit par seconde

Tbit/s Térahit par seconde

UNITÉ GÉOMÉTRIQUE

km kilomètre

nm nanomètre

UNITÉ DE DÉBIT

dB décibel

UNITÉ DE PUISSANCE

W Watt

λ longueur d'onde

INTRODUCTION

Le développement des nouvelles technologies et applications (5G, vidéo, infonuagique, etc.) ainsi que l'implémentation des centres de données ont conduit à une augmentation constante du trafic Internet au fil des ans. Pour faire face à cela, les opérateurs de réseaux de télécommunication se voient contraints de déployer des systèmes de transmission optiques. Ces derniers ont pour but principal de répondre aux besoins grandissants du volume de données échangées dans les réseaux. Ils offrent des débits beaucoup plus rapides (de 100 Gbit/s à des Tbit/s) sur de plus longues distances (pouvant aller jusqu'à 10,000 km sans régénération), permettant ainsi d'avoir des réseaux avec des débits plus élevés (Mukherjee, 2020).

0.1 Problématique

Les signaux optiques générés par des systèmes de transmission de capacité de plus en plus élevée, suite à l'évolution continue du trafic, sont sujets à plusieurs limitations ayant un impact important sur la qualité de transmission. Parmi ces limitations, l'on trouve notamment :

- les pertes optiques causées par l'atténuation dans la fibre optique et les pertes dues aux connecteurs, épissures et composants,
- la dispersion entraînant un élargissement temporel du signal optique,
- les effets non linéaires dans la fibre optique.

En outre, les applications utilisant les systèmes de transmission déployés requièrent à la fois une fiabilité et une qualité de transmission accrues du réseau, ainsi qu'une hétérogénéité et une flexibilité élevée de celui-ci. Il devient donc important de minimiser l'impact des limitations observées dans le réseau. L'une des solutions serait d'implémenter des outils proactifs de conception, de gestion et de contrôle des réseaux optiques. Cependant, les approches courantes pour l'implémentation des systèmes de contrôle et de gestion des réseaux n'exploitent pas pleinement les métriques de performance du réseau afin de fournir de la proactivité dans la gestion des ressources du réseau. Ces approches se basent sur des valeurs de marge système

élevées (Panayiotou et al., 2019a; T. Panayiotou, 2017) afin d'éviter les défaillances et les problèmes de qualité de transmission dans le réseau. Cela entraîne une surutilisation des ressources du réseau. D'autres approches de contrôle et de gestion du réseau utilisent des modèles analytiques (Azodolmolky et al., 2011; Azodolmolky et al., 2009) requérant des temps de calcul élevés et des processeurs puissants (Amirabadi, 2019).

Dans ce contexte, la notion de cognition, introduite et déployée dans le monde radio, se veut une solution dans le domaine optique. Elle apporte de l'intelligence dans la gestion et le contrôle du réseau. Elle permet également d'avoir des méthodes de gestion autonome des performances du réseau et de minimiser les interventions humaines dans les réseaux optiques (Borkowski et al., 2015). Un système cognitif est défini comme étant un système intelligent qui, en se basant sur son environnement extérieur et sur son historique, a la capacité de prendre des décisions afin d'ajuster ses paramètres de transmission par rapport à son état actuel (Mahmoud, 2007). En d'autres termes, avec l'observation, l'analyse et l'apprentissage de son environnement interne et externe (taux d'erreur binaire (*Bit Error Rate*, BER), ou rapport signal-sur-bruit optique (*optical signal-to-noise ratio*, OSNR), ou température, etc.), ce dernier permet d'optimiser le fonctionnement du réseau (sélection des longueurs d'onde ou du format de modulation le plus approprié des canaux, sélection du meilleur chemin, etc.). À cet effet, les algorithmes d'apprentissage machine sont proposés dans l'implémentation des systèmes optiques cognitifs.

0.2 Contributions

L'objectif principal de cette thèse est de proposer de nouveaux modèles basés sur des algorithmes d'apprentissage machine dans le but d'estimer la qualité de transmission des connexions optiques et de prédire leur performance une fois déployées dans un réseau. L'un des aspects originaux de ces travaux est le fait que ces modèles ont été construits et validés avec des données de performance réelles recueillies dans des réseaux de production. Cela permet ainsi à chacun de ces modèles d'être plus représentatifs du contexte réel auquel font

face les opérateurs de télécommunication comparativement aux travaux antérieurs réalisés avec des données synthétiques. Pour ce faire, plusieurs contributions secondaires ont pu être établies.

0.2.1 Estimateur de performance d'une connexion optique

La première contribution est le développement d'un estimateur de performance d'une connexion optique. En exploitant des algorithmes d'apprentissage machine existants, il est question d'implémenter un module de classification permettant de caractériser le niveau de qualité (BER) d'une connexion optique. Notons que cet estimateur est testé à partir des données de BER recueillies sur une connexion optique déployée dans un réseau de production. En outre, l'une des particularités est également l'implémentation dans le module de classification d'une ingénierie d'attributs ayant pour but de sélectionner les données d'entrée du module les plus pertinentes pour la classification.

0.2.2 Prédicteur de performance d'une connexion optique

La deuxième contribution est le développement d'un prédicteur de performance dans le but de prédire les valeurs de performance d'une connexion optique et de détecter les dégradations de performance. Le prédicteur de performance est implémenté et testé à partir des paramètres de performance recueillis sur une connexion optique (BER, puissance du canal, etc.) d'un réseau de production. Une étude d'ingénierie est également effectuée sur ces paramètres afin de sélectionner les données d'entrée les plus pertinentes pour le modèle. Ainsi, le prédicteur de performance utilise des données de BER seulement ou des données supplémentaires en plus du BER. En outre, notons que c'est la première fois qu'un prédicteur de performance multivarié a été proposé pour prédire la performance d'une connexion optique sur un horizon de quelques jours.

0.2.3 Détecteur des anomalies de performance d'une connexion optique

La troisième contribution est le développement d'un détecteur d'anomalies de performance de connexion optique. Il permet de définir les anomalies en utilisant une méthode différente de l'approche basée sur des seuils fixés par les opérateurs de télécommunication et de détecter les anomalies survenant dans une connexion optique. Tout comme les outils implémentés dans les contributions précédentes, le détecteur d'anomalies est implémenté en utilisant des algorithmes d'apprentissage machine existants et est testé avec des données de performance recueillies sur des connexions optiques d'un réseau de production.

0.2.4 Contributions scientifiques

Les résultats de cette thèse ont également donné lieu à des publications dans des articles de revue et des présentations à des conférences :

- S. Allogba, S. Aladin and C. Tremblay. 2021. « Multivariate Machine Learning Models for Short-Term Lightpath Performance Forecast », in Journal of Lightwave Technology, doi: 10.1109/JLT.2021.3110513.
- S. Allogba, B. L. Yameogo and C. Tremblay. 2021. « Extraction and Early Detection of Anomalies using Machine Learning Models », in Journal of Lightwave Technology (version révisée *soumise le 20 octobre 2021*).
- S. Aladin, A. V. S. Tran, S. Allogba et C. Tremblay. 2020. « Quality of Transmission Estimation and Short-Term Performance Forecast of Lightpaths ». Journal of Lightwave Technology, vol. 38, no 10, p. 2807-2814.
- S. Aladin, S. Allogba, A. V. S. Tran and C. Tremblay, "Recurrent Neural Networks for Short-Term Forecast of Lightpath Performance," 2020 Optical Fiber Communications Conference and Exhibition (OFC), San Diego, CA, USA, 2020, Paper W2A.24.
- C. Tremblay, S. Allogba and S. Aladin, « Quality of transmission estimation and performance prediction of lightpaths using machine learning » 45th European Conference on Optical Communication (ECOC 2019), Dublin, Ireland, 2019, pp. 1-3.

- S. Allogba et C. Tremblay. 2018. « K-Nearest Neighbors Classifier for Field Bit Error Rate Data ». In 2018 Asia Communications and Photonics Conference (ACP). (26-29 Oct. 2018), p. 1-3 (*Best Paper Award Competition*).
- S. Allogba et C. Tremblay. June 5-7, 2018. « Statistical analysis of performance monitoring data collected in an optical link ». In 20th Photonics North Conf. (Montreal, Canada, June 5-7, 2018).
- S. Allogba et C. Tremblay. June 5-7, 2017. « Prediction of BER Trends in an Optical Link Using Machine Learning ». In 19th Photonics North Conf. (Ottawa, Canada, June 5-7, 2017).

0.3 Organisation

Cette thèse est divisée en cinq chapitres. Dans le premier chapitre, une présentation des réseaux optiques ainsi que des algorithmes d'apprentissage machine est faite. Il a pour but de fournir un aperçu des technologies étudiées et de définir les problèmes de performance rencontrés dans les réseaux. On y présente également les différentes méthodes cognitives utilisées dans les réseaux optiques, ainsi qu'un récapitulatif des avantages et des faiblesses de chacune de ces méthodes proposées.

Le deuxième chapitre de la thèse présente la méthode de classification des données de performance (BER) d'une connexion optique. Celui-ci sera subdivisé en deux parties :

- la proposition d'un premier outil de classification des données de BER recueillies dans un réseau de production incluant les étapes effectuées pour la classification des données de performances du réseau et les résultats de simulation obtenus avec l'algorithme de classification ;
- l'analyse de l'impact des attributs permettant l'estimation de la qualité de transmission du signal optique.

Les chapitres suivants de la thèse sont consacrés au problème de prédiction du rapport signal-sur-bruit (*signal-to-noise ratio*, SNR) en se basant sur des données historiques de performance collectées dans des liaisons optiques d'un réseau de production. Des modèles univariés et multivariés de prédiction du SNR sont proposés dans le troisième chapitre et le quatrième chapitre de la thèse respectivement. Chacun de ces chapitres décrit les étapes d'implémentation des méthodes de prédiction ainsi que les résultats obtenus. Le quatrième chapitre présente également une piste d'optimisation des méthodes de prédiction via l'apprentissage par transfert.

Le cinquième chapitre présente les méthodes utilisées pour la détection anticipée des anomalies observées sur les données de performance collectées dans des liaisons optiques. Une description de ces données de performance est effectuée afin de définir la notion d'anomalies et d'identifier les principaux attributs ayant un impact sur ces anomalies. Ce chapitre décrit également les méthodes de sélection des attributs pertinents pour la détection des anomalies.

CHAPITRE 1

REVUE DE LA LITTÉRATURE

1.1 Introduction

Afin de répondre à l'évolution continue des nouvelles applications technologiques, les systèmes optiques déployés se retrouvent contraints à transporter de plus en plus de volume de trafic sur de longues distances. À cet effet, les pannes et événements survenant sur le réseau ont de plus en plus d'influence sur les performances du réseau.

Les systèmes de contrôle et de gestion dynamique du réseau se présentent donc comme des solutions adéquates afin de minimiser l'impact des pannes sur la qualité de service. Ces systèmes de contrôle et de gestion dynamique permettraient d'exploiter les paramètres de performance du réseau dans le but de détecter ou de prédire des anomalies pouvant survenir sur le réseau. Aussi, les algorithmes d'apprentissage automatique ou apprentissage machine, bien connus dans le domaine des réseaux sans fil, peuvent-ils être utilisés dans ces systèmes pour permettre au réseau d'apprendre des événements passés afin de prendre des décisions proactives (Thomas et al., 2006).

Ce chapitre est subdivisé en trois sections. Les deux premières sections présentent une description des réseaux optiques et de certaines méthodes d'apprentissage automatique dans le but de mieux saisir leur fonctionnement et leurs caractéristiques. La dernière section fait l'état de l'art des méthodes d'apprentissage automatique dans le domaine des communications optiques.

1.2 Réseaux optiques : présentation générale

1.2.1 Description générale

Déployée partout dans le monde de la télécommunication pour répondre aux faiblesses des réseaux traditionnels et pour permettre des télécommunications à l'échelle locale, nationale et internationale, la fibre optique présente des avantages indéniables sur différents aspects tels que (Mukherjee, 2020) :

- la transmission de très grandes capacités de données sur de très longues distances grâce à la bande passante (plus de 50 Tbit/s) élevée et à la faible atténuation du signal transporté dans la fibre optique,
- l'absence d'interférence entre les fibres optiques d'un même câble, l'immunité aux interférences électromagnétiques et la faible distorsion du signal,
- le potentiel de sécurité inhérent à la fibre optique,
- la faible consommation électrique et le coût réduit des équipements de transmission optique comparativement aux autres technologies.

En outre, comme tout réseau de télécommunication, les réseaux optiques sont constitués de plusieurs nœuds interconnectés entre eux à travers un canal de transmission. Dans le contexte des réseaux optiques, ce canal de transmission est la fibre optique.

De plus, le réseau optique peut ne pas être composé que d'équipements tout optiques. En effet, hormis la fibre optique utilisée comme canal de transmission, les équipements de génération et de détection du signal ainsi que les composants d'interconnexion peuvent être tout optiques, électroniques ou encore hybrides (à la fois électroniques et optiques).

La Figure 1.1 présente un exemple de réseau optique, constitué de plusieurs chemins ou liaisons optiques reliant les différents terminaux (transmetteur et/ou récepteur) du réseau. Ces

chemins optiques sont composés de plusieurs éléments permettant de transporter le trafic d'un terminal à un autre. Ainsi, les éléments constitutifs d'un réseau optique sont :

- les terminaux : les terminaux sont les points source et destination de trafic. Ils comprennent les transmetteurs/récepteurs qui assurent à la fois la conversion du signal électrique en signal optique (transmetteur) et la conversion du signal optique en signal électrique (récepteur). Chaque terminal inclut également des multiplexeurs (MUX) et des démultiplexeurs (DEMUX) permettant respectivement de regrouper ou de démultiplexer les signaux optiques (ou canaux) WDM (*Wavelength Division Multiplexing*) dans la même fibre optique. Chaque canal WDM possède une longueur d'onde spécifique. Les terminaux peuvent également contenir des équipements de monitoring pour enregistrer les paramètres de performance du réseau ainsi que des systèmes de gestion et de contrôle du réseau ;
- les nœuds intermédiaires entre les points source et destination du trafic dans lesquels peuvent se trouver des multiplexeurs optiques configurables d'injection-extraction de longueur d'onde (*Reconfigurable Optical Add-Drop Multiplexer, ROADM*) et commutateurs sélectifs en longueur d'onde (*Wavelength Selective Switch, WSS*) permettant d'extraire ou d'ajouter des longueurs d'onde aux sites intermédiaires. Ces nœuds peuvent également contenir des équipements de monitoring pour enregistrer les paramètres de performance du réseau ;
- les amplificateurs optiques pour compenser les pertes occasionnées par l'atténuation dans la fibre optique et les pertes d'insertion des composants ;
- les sites de régénération permettant de régénérer le signal optique via la conversion du signal optique en signal électrique et sa reconversion en signal optique.

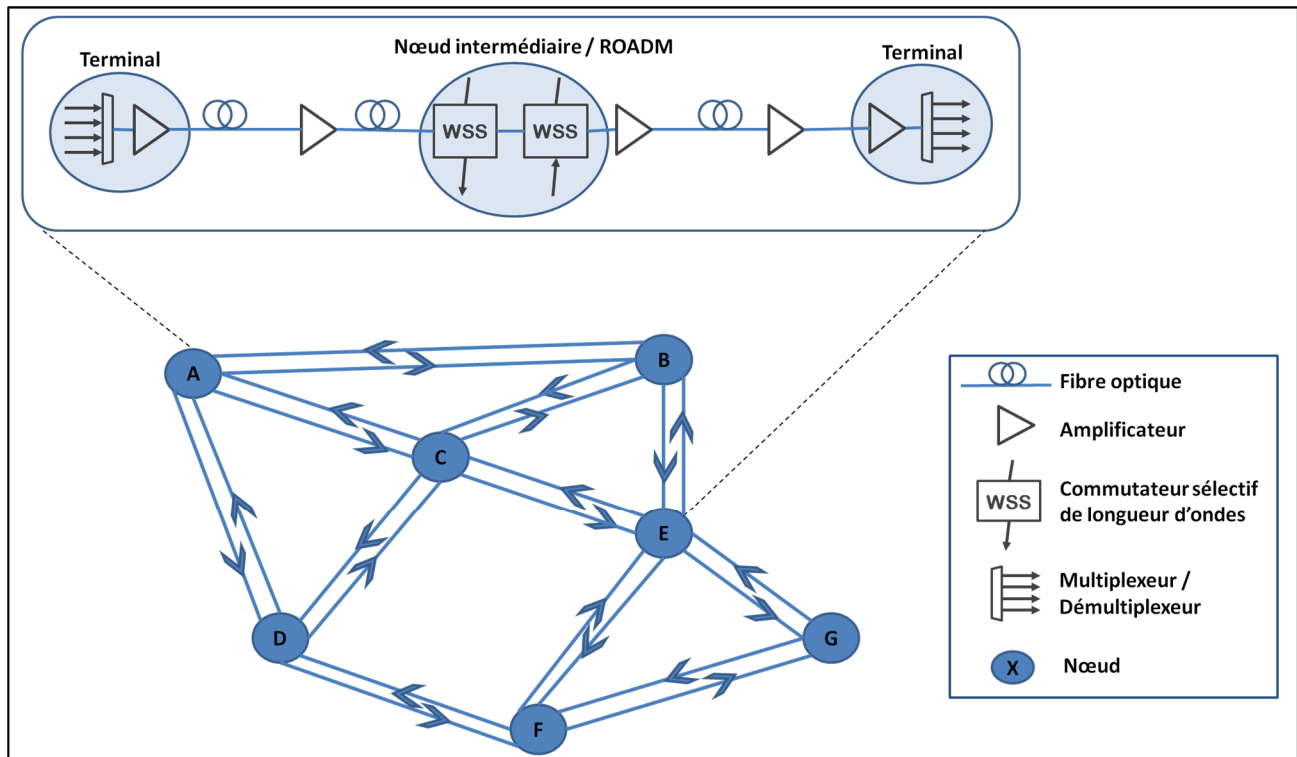


Figure 1.1 Schéma général d'un réseau optique

Le signal circulant dans le réseau subit plusieurs perturbations qui réduisent les performances du réseau. À cet effet, le prochain paragraphe présente les effets de dégradation du signal.

1.2.2 Effets de dégradation du signal optique

Les effets de dégradation du signal optique peuvent être catégorisés en deux groupes : les effets linéaires regroupant l'atténuation, la dispersion chromatique ainsi que la dispersion des modes de polarisation et les effets non linéaires résultant de la propagation de l'onde électromagnétique lumineuse dans la fibre optique (Simmons, 2014).

1.2.2.1 Atténuation

L'atténuation a un impact important sur la transmission d'un signal dans une fibre optique. Fonction de la distance entre le transmetteur et le récepteur, l'atténuation, causée par la diffusion et l'absorption du signal, se définit par l'équation (1.1) (Mukherjee, 2006). Elle se traduit par la diminution de la puissance optique du signal pendant sa propagation dans la fibre optique.

$$\alpha = \frac{10}{z} \log \left(\frac{P_0}{P_z} \right) \quad (1.1)$$

où :

- α représente le coefficient d'atténuation exprimé en dB/km ;
- P_0 et P_z représentent respectivement la puissance du signal à l'origine ($z = 0$) et à la distance z , exprimée en W ;
- z représente la distance parcourue par le signal dans la fibre optique, exprimée en km.

À l'atténuation dans les fibres optiques de silice s'ajoutent les pertes optiques associées aux connecteurs et épissures reliant les segments de câble optique ainsi que les pertes d'insertion des composants déployés dans les liaisons optiques. Les pertes d'atténuation et d'insertion peuvent être compensées par le déploiement de sites d'amplification optique dans la connexion optique.

1.2.2.2 Dispersion chromatique

Le phénomène de dispersion chromatique dans la fibre optique engendre un étalement temporel des impulsions optiques lors de leur propagation dans la fibre optique. Le phénomène est causé par le fait que l'indice de réfraction du matériau constituant la fibre optique, correspondant au ratio entre la vitesse de la lumière dans le vide et la vitesse de la lumière dans le milieu, est différent d'une longueur d'onde à une autre. Ainsi, lors de la transmission du

signal optique dans la fibre optique, les canaux WDM se propagent chacun avec une vitesse différente, avec pour résultat un étalement temporel des impulsions. La dépendance de l'indice de réfraction avec la longueur d'onde fait en sorte que la largeur spectrale de la source lumineuse a un impact sur l'ampleur de l'effet de dispersion tout comme la distance. L'effet de dispersion chromatique est représenté par l'équation (1.2) (Keiser, 2010) :

$$\Delta t = D\sigma_{\lambda}L \quad (1.2)$$

où :

- Δt représente la dispersion chromatique, exprimée habituellement en ps étant donné les ordres de grandeur en jeu ;
- D représente le coefficient de dispersion de la fibre optique, exprimé en ps/(km.nm). Le coefficient de dispersion (de l'ordre de 19 ps/nm.km à 1550 nm pour une fibre optique de type ITU-T G.652) dépend du type de fibre optique et de la longueur d'onde, spécifiés par les manufacturiers ;
- σ_{λ} représente la largeur spectrale de la source à la longueur d'onde λ , exprimée en nm ;
- L représente la longueur de la connexion optique, exprimée en km.

1.2.2.3 Dispersion des modes de polarisation

Tout comme la dispersion chromatique, la dispersion des modes de polarisation (*Polarization Mode Dispersion, PMD*) est un effet d'étalement temporel des impulsions dans la fibre optique. Elle est la résultante du délai de propagation de chacun des deux modes de polarisation du signal dans la fibre optique (*Differential Group Delay, DGD*), tel que présentée par l'équation (1.3), et est causée par la biréfringence dans la fibre optique. La biréfringence de la fibre optique est un phénomène dû aux caractéristiques intrinsèques de la fibre optique (comme l'ellipticité du cœur) ou aux contraintes extérieures (comme des torsions de la fibre optique) (Keiser, 2010).

$$\Delta\tau = D_{PMD}\sqrt{L} \quad (1.3)$$

où :

- $\Delta\tau$ représente l'étalement temporel des impulsions causé par la DGD, en ps ;
- D_{PMD} représente le coefficient de PMD, une spécification de la fibre optique exprimée en $ps/\sqrt{km}$ qui varie entre 0,1 et 1 $ps/\sqrt{km}$ de manière typique ;
- L représente la longueur de la liaison en km.

1.2.2.4 Effets non linéaires

Les effets non linéaires de dégradation du signal sont causés par les phénomènes de diffusion inélastique et Kerr dans la fibre optique. À cet effet, plus la puissance du signal augmente, plus les effets non linéaires de dégradation deviennent importants (Simmons, 2014). Les effets non linéaires observés sont les suivants :

- les effets de modulation de phase, à savoir la modulation de phase croisée (*Cross-phase modulation, XPM*) et l'automodulation de phase (*Self-Phase Modulation, SPM*) qui entraînent des effets de déphasage du signal qui impactent chaque canal WDM et des effets d'interférence sur les autres canaux WDM ;
- le mélange à 4 ondes (*Four Wave Mixing, FWM*), semblable aux effets d'interférence inter-symbole dans le domaine électrique, fait référence à l'apparition de signaux parasites générés par l'interaction des canaux WDM ;
- les effets de diffusion stimulée Brillouin (*Stimulated Brillouin Scattering, SBS*) et Raman (*Stimulated Raman Scattering, SRS*) causée par l'interaction du signal avec des atomes de silice.

L'établissement d'une connexion optique dans un réseau ainsi que la gestion et le contrôle de sa qualité de transmission dépendent de plusieurs paramètres influencés par les effets de dégradation précités. La prochaine section décrit les principaux paramètres observés pour caractériser la performance des connexions optiques déployées dans un réseau.

1.2.3 Paramètres de performance

Les principaux paramètres de performance utilisés pour décrire la qualité de transmission d'un signal optique sont le taux d'erreur binaire (*Bit Error Rate, BER*), le rapport signal-sur-bruit (SNR), le SNR optique (*Optical signal-to-noise ratio, OSNR*) et le facteur Q. Le facteur Q et le SNR étant des métriques de performance dérivées du BER, seul le taux d'erreur binaire et l'OSNR sont présentés dans cette section.

1.2.3.1 Taux d'erreur binaire

Le taux d'erreur binaire traduit la sensibilité ou la qualité du récepteur dans la connexion optique et se définit comme étant le niveau d'erreur binaire acceptable sur une liaison à la suite des effets de dégradation. Ainsi, meilleure est la qualité du signal, plus faible est le BER (Chan, 2010).

En outre, afin de respecter les exigences de qualité de transmission d'une connexion optique, les opérateurs de télécommunication fixent des niveaux de pré-FEC BER très faible (10^{-9} ou 10^{-12} , par exemple). Les codes correcteurs d'erreur (*Forward Error Correction, FEC*) ont été implémentés dans le but de détecter et de corriger les erreurs au récepteur et de permettre d'opérer les connexions à des valeurs minimums de BER (post-FEC BER) jusqu'à 10^{-3} . Le pré-FEC BER est un paramètre primordial pour déterminer les performances d'un signal optique.

1.2.3.2 Rapport signal-sur-bruit optique

Le rapport signal-sur-bruit optique (OSNR) représente un critère de performance d'un signal optique dans les liaisons optiques (Chan, 2010). Plus l'OSNR est élevé, plus la qualité du signal est élevée. L'OSNR permet d'estimer la qualité du signal d'une connexion optique et est déterminé par le rapport de la puissance du signal sur celle du bruit (la principale source de

bruit étant celui généré par les amplificateurs optiques déployés dans la liaison) tel que décrit dans l'équation (1.4) (Mukherjee, 2020) :

$$OSNR = 10 \log \left(\frac{P}{2 B_{ref} N_{ase}} \right) \quad (1.4)$$

où :

- P représente la puissance optique du signal, exprimée en W ;
- B_{ref} représente la bande passante de référence utilisée pour la mesure du bruit, exprimée en GHz ;
- N_{ase} représente la densité spectrale de la puissance de bruit d'un amplificateur optique, exprimé en W.

Les transmetteurs et récepteurs optiques cohérents permettent de collecter les données de performance du réseau via les modules de traitement de signal (*Digital Signal Processing, DSP*) intégrés dans ces composants (Roberts et al., 2017; Sambo et al., 2015). Des puissancemètres optiques (*optical power meters*) sont également déployés le long des liaisons afin de monitorer les niveaux de puissance du signal.

Les opérateurs réseau établissent les liaisons optiques avec des marges système élevées afin de prévenir le plus possible les pannes et les fautes. Les marges systèmes sont définies comme étant l'écart, en dB, entre le niveau de performance requis au récepteur pour obtenir le BER minimum requis et le niveau de BER effectif sur le terrain. Développer des systèmes de contrôle proactifs utilisant les données collectées dans les réseaux pour anticiper les pannes et les dégradations du signal permettraient d'éviter les pannes et la surutilisation des ressources du réseau lors du routage ou de l'assignation des longueurs d'onde dans le réseau. Les techniques d'apprentissage machine utilisant des données passées afin de poser des décisions (prédiction, contrôle) pourraient s'avérer être des éléments essentiels de ces nouveaux

systèmes proactifs de contrôle. La prochaine section fait une présentation générale des méthodes d'apprentissage machine.

1.3 Apprentissage machine

1.3.1 Présentation générale

L'apprentissage machine est une technique utilisant des algorithmes pour permettre à un système de s'adapter ou de modifier son fonctionnement en fonction des observations et des analyses continues de données issues de capteurs présents dans le système (Marsland, 2014; Witten, Frank et Hall, 2011). Ces algorithmes apportent de la cognition aux systèmes dans lesquels ils sont implémentés, leur permettant d'interagir avec leur environnement par l'apprentissage, l'identification et la prédiction des différents états du système (Arslan, 2007).

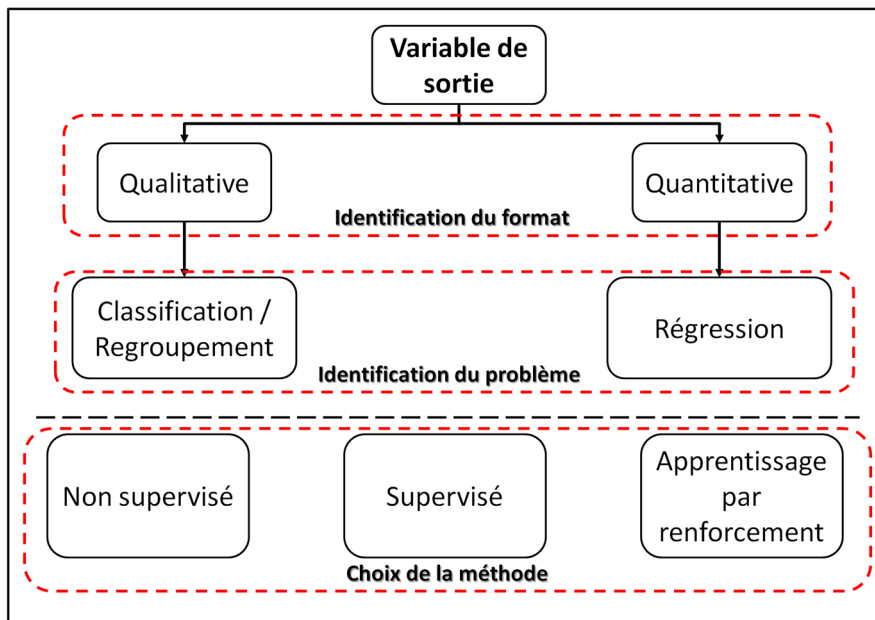


Figure 1.2 Étapes d'identification de la méthode d'apprentissage machine

Comme présenté dans la Figure 1.2, le choix de la méthode d'apprentissage machine se fait en trois étapes. La première étape consiste en l'identification du format de la variable de sortie. En effet, la variable de sortie oriente le choix du problème d'apprentissage machine. Celle-ci est définie à partir de l'objectif fixé par l'opérateur (exemple : prédiction de la valeur d'un paramètre de performance, détection d'une panne sur le réseau, estimation de l'état de la qualité de la connexion optique ...). Il existe deux types de données de sortie : les données quantitatives et les données qualitatives. Lorsque l'objectif de l'opérateur est la prédiction d'une valeur (exemple : prédiction de la valeur du BER dans le cas du réseau optique), les données sont de type quantitatif. La variable de sortie est donc numérique. Par ailleurs, lorsque l'objectif de l'opérateur est de prédire la catégorie d'une donnée ou d'un paramètre (exemple : prédiction de la qualité (bonne ou mauvaise) d'une connexion optique), les données sont de type qualitatif. Dans ce cas, la variable de sortie est de type catégorique (exemple : 0 ou 1, bon ou mauvais).

La seconde étape consiste à identifier le type de problème d'apprentissage machine en se basant sur le format de la variable de sortie définie. En effet, si la variable de sortie est de type qualitatif, le problème d'apprentissage machine est un problème de classification ou de regroupement ou de partitionnement (*clustering*). Dans le cas contraire, lorsque la variable de sortie est de type quantitatif, le problème d'apprentissage machine est un problème de régression.

La dernière étape consiste au choix de la méthode. Celle-ci est fonction de la connaissance des informations contenues dans l'objectif de départ (type de données, connaissance des classes ...). Il existe trois grands groupes de méthodes d'apprentissage machine : les méthodes supervisées, les méthodes non supervisées et les méthodes d'apprentissage par renforcement.

1.3.2 Méthodologie générale

1.3.2.1 Description des termes

Afin de comprendre le fonctionnement des algorithmes d'apprentissage machine, il convient de décrire les principaux termes utilisés en apprentissage machine, à savoir : la base de données, les étiquettes ou labels et les attributs ou caractéristiques.

- Base de données : la base de données comprend les informations d'entrée et de sortie fournies par l'opérateur (exemple : valeurs temporelles de BER). Lors de l'implémentation des algorithmes d'apprentissage machine, elle est souvent divisée en trois groupes. Chacun des groupes correspond aux différentes phases d'implémentation du modèle. La première phase est la phase de validation (utilisant les données de validation). En effet, certains modèles d'apprentissage machine sont caractérisés par plusieurs hyper-paramètres représentant leurs paramètres intrinsèques. Les hyper-paramètres optimaux sont déterminés en évaluant le modèle avec le jeu de données de validation. Par la suite, la seconde phase est la phase d'apprentissage (utilisant les données d'apprentissage). Le modèle formé dans la phase de validation est entraîné en utilisant ce jeu de données afin d'apprendre les caractéristiques du système. Enfin, la dernière phase est la phase de test (utilisant les données de test). C'est le lieu d'évaluer les performances du modèle implémenté en comparant la sortie du système avec la sortie réelle de l'ensemble de test. L'implémentation de certains modèles ne nécessite pas les données de validation. Les hyper-paramètres sont alors définis selon l'opérateur et le modèle est entraîné en utilisant les données d'apprentissage. Les données de test sont utilisées pour évaluer les performances du modèle. Par ailleurs, en remplacement de la division de la base de données selon un ratio fixe en deux ou trois groupes (validation, apprentissage, test), une autre technique, appelée validation croisée, consiste à diviser la base de données en plusieurs sous-ensembles rotatifs et à évaluer les performances du modèle pour chacun des sous-

ensembles. Le modèle final est évalué en se basant sur la performance moyenne obtenue à l'issue de la compilation du modèle pour chacun des sous-ensembles.

- Étiquettes (*labels*) : les étiquettes correspondent à la valeur de sortie désirée. Elles sont généralement définies dans le cas de classification lorsque l'objectif est de prédire une variable qualitative. C'est le cas, par exemple, de la prédiction de la qualité de transmission d'une connexion optique dans une liaison. L'étiquette pouvant être attribuée est « bonne » ou « mauvaise ».
- Attributs ou caractéristiques (*features*) : ils servent à décrire la donnée à prédire. Ils peuvent être de format numérique ou catégorique. Dans le cas de la prédiction du BER, par exemple, les attributs pouvant être utilisés sont : la température, la période de la journée, les paramètres du réseau, etc. Le choix des attributs est important, car ceux-ci ajoutent de l'information au modèle à implémenter. Cependant, si les attributs ajoutés ne sont pas en lien avec la donnée à prédire, ils peuvent dégrader la performance du modèle. Des critères de sélection sont mis en place afin d'évaluer le lien entre chaque attribut et la donnée étudiée (Chandrashekar et Sahin, 2014).
- Métriques : elles permettent d'évaluer les performances des algorithmes. Il n'existe pas de métriques spécifiques pour l'évaluation des algorithmes. En effet, les métriques à choisir peuvent varier en fonction du problème (classification, partitionnement ou régression) et du type de données. Dans le cadre d'une régression, les principales métriques sont l'erreur quadratique moyenne (*Root Mean Square Error, RMSE*), l'erreur absolue moyenne (*Mean Absolute Error, MAE*), l'erreur absolue maximale (*Absolute Maximum error, AME*) et la valeur moyenne quadratique ou R-carré (R^2). Dans le cadre d'une classification, les métriques utilisées sont entre autres la justesse, la précision, la courbe ROC (*Receiver Operating Characteristic*) et la matrice de confusion. À partir de la matrice de confusion, d'autres métriques, telles que le taux de faux positif, le taux de faux négatif, etc., peuvent être également définies.

1.3.2.2 Présentation des méthodes d'apprentissage machine

Les algorithmes d'apprentissage machine sont catégorisés en trois grands groupes : l'apprentissage supervisé, l'apprentissage non supervisé et l'apprentissage par renforcement. Dans les travaux présentés dans cette thèse, les méthodes utilisées font partie de la famille de l'apprentissage supervisé. Ainsi, cette section décrit l'apprentissage supervisé et fait également une présentation générale de l'apprentissage non supervisé et l'apprentissage par renforcement.

L'apprentissage supervisé regroupe les algorithmes de régression et les algorithmes de classification. Dans ce groupe, toutes les informations de la base de données incluant les attributs et les labels ou les valeurs de sorties sont connues (Gu, Yang et Ji, 2020). En d'autres termes, après la répartition de la base de données en données d'apprentissage et de test, le but des algorithmes d'apprentissage supervisé est de prédire une donnée $\hat{Y}$ ($\hat{Y}$ pouvant être une valeur numérique ou une étiquette). Cela se fait en se basant sur l'association des éléments connus de la base d'apprentissage (X_i, Y_i) , avec X correspondant aux attributs et Y à la donnée de sortie (ou étiquette) associée (Mohri, Rostamizadeh et Talwalkar, 2012; Singh, Thakur et Sharma, 2016). Les algorithmes d'apprentissage supervisé sont implémentés en quatre étapes (Kotsiantis, 2007; Singh, Thakur et Sharma, 2016). La première étape est la collecte des données et des attributs associés. La seconde étape est le prétraitement des données (traitement des données manquantes, répartition de la base de données en données d'apprentissage, validation et test ...). La troisième étape est le traitement des attributs (réduction du nombre d'attributs, identification des attributs pertinents ...). La dernière étape est le choix de l'algorithme d'apprentissage machine. Celui-ci est fonction du problème visé (classification ou régression) et des performances désirées (temps d'exécution, temps d'apprentissage, précision ...). En outre, les algorithmes d'apprentissage supervisés sont classés en deux groupes (Francesco Musumeci et al., 2018) : les modèles paramétriques et les modèles non paramétriques. Parmi les modèles paramétriques se trouvent les réseaux de neurones (*Neural Networks, NN*) et les modèles de régression logistique. Dans ces modèles, le nombre de paramètres est fixe. Ils sont estimés dans la phase d'apprentissage et de validation. Par ailleurs,

les modèles non paramétriques ont un nombre de paramètres en fonction de l'ensemble d'apprentissage. Les algorithmes les plus utilisés sont les machines à vecteur support (*Support Vector Machine, SVM*), les arbres de décisions et les algorithmes du plus proche voisin (*K Nearest Neighbor, K-NN*).

Contrairement aux modèles d'apprentissage supervisé, les labels des modèles d'apprentissage non supervisés ne sont pas définis (Gu, Yang et Ji, 2020; Mohri, Rostamizadeh et Talwalkar, 2012). En d'autres termes, le but des modèles d'apprentissage non supervisé est de prédire ou catégoriser une donnée $\hat{Y}$ en ne se basant que sur les similarités présentes avec les éléments d'entrée X de la base d'apprentissage, les labels Y étant inconnus. Ces modèles sont utilisés dans trois catégories, à savoir le partitionnement (*clustering*), l'estimation de la densité et la réduction de dimensionnalité. Les deux premières catégories sont utilisées pour déterminer la classe d'une donnée. En effet, les algorithmes de partitionnement, équivalents des algorithmes de classification en apprentissage supervisé, permettent de prédire la classe d'une donnée en fonction des similarités observées dans l'ensemble d'apprentissage. L'algorithme le plus utilisé est le partitionnement en k-moyenne (*K-mean*). Les algorithmes d'estimation de densité prédisent la classe d'une donnée à partir de la détermination de sa distribution ou de l'estimation des paramètres de sa distribution. L'algorithme utilisé dans cette catégorie est l'algorithme Espérance Maximisation (*Expectation Maximization, EM*). Par ailleurs, les algorithmes de réduction de dimensionnalité sont utilisés lors du prétraitement des données pour la compression des données ou la réduction des attributs. L'algorithme principal utilisé dans cette catégorie est l'analyse en composantes principales ou ACP (*Principal Component Analysis, PCA*). Les réseaux de neurones peuvent également être utilisés pour la réduction de dimensionnalité (Bishop, 2006).

Les algorithmes d'apprentissage par renforcement sont basés sur l'application de récompense pour chaque action posée et utilisent les processus de décision markovien (Francesco Musumeci et al., 2018; Gao et al., 2020). L'objectif des algorithmes d'apprentissage par renforcement est alors de déterminer la donnée de sortie Y qui donnera la meilleure

récompense. À l'inverse des algorithmes d'apprentissage supervisé, les algorithmes d'apprentissage par renforcement ne dépendent pas des étiquettes définies par l'expert. Ils sont fonction de l'expérience du modèle définie à partir de son environnement et des récompenses et/ou pénalités attribuées à chaque prédiction effectuée. L'algorithme le plus couramment utilisé dans l'apprentissage par renforcement est le *Q-learning*.

Les algorithmes d'apprentissage machine ont progressivement été insérés dans le monde de la communication optique afin d'améliorer les performances du réseau. Ainsi, la prochaine section fait une revue de littérature sur les différentes utilisations des algorithmes d'apprentissage machine dans le domaine des communications optiques.

1.4 Apprentissage machine dans les télécommunications optiques

En plus de requérir des systèmes complexes avec des temps de calcul élevés, les outils traditionnels de gestion et de contrôle des réseaux optiques n'utilisent pas de façon optimale les ressources du réseau et gèrent les pannes et défaillances en mode réactif (Gu, Yang et Ji, 2020). Ainsi, l'utilisation des méthodes d'apprentissage machine pourrait minimiser les interventions humaines, de par leur apport d'intelligence, d'automatisation et de prédiction dans le réseau. Ces méthodes permettraient également d'améliorer la flexibilité de l'utilisation des ressources du réseau, à travers l'utilisation des informations internes et externes du réseau (performances collectées, état du réseau ...). Plusieurs études ont recensé les différentes applications des méthodes d'apprentissage machine dans le domaine des communications optiques les regroupant dans trois domaines : le suivi des performances, la détection des pannes, l'estimation de la qualité de transmission (Gu, Yang et Ji, 2020; Mata et al., 2018b; Musumeci et al., 2019b). Cette section sera axée sur l'estimation de la qualité de transmission ainsi que la détection des pannes.

1.4.1 Estimation de la qualité de transmission

L'implémentation des méthodes d'apprentissage machine dans l'estimation de la qualité de transmission répond aux problèmes de gaspillage des ressources lors de l'établissement de nouvelles connexions (Gu, Yang et Ji, 2020). En se basant sur les paramètres du réseau, elles permettent donc de prédire en temps réel l'état d'une nouvelle connexion optique à établir sans avoir à passer par des calculs analytiques complexes dans un contexte où tous les paramètres de calcul ne sont pas connus.

L'estimation de la qualité de transmission est un problème de classification. En effet, le but est de prédire la qualité de transmission (« bonne » ou « mauvaise ») de la nouvelle connexion optique à établir. Dans ce contexte, L. Barletta (March 2017) et Cristina Rottondi (2018) définissent un niveau de BER à partir duquel la connexion à établir peut être classée comme acceptable ou non. Par la suite, en se basant sur les paramètres physiques du réseau et du signal comme attributs, à savoir : la longueur du chemin optique, le nombre de liaisons sur le chemin, la longueur maximale des liaisons, le volume du trafic et le format de modulation, ils utilisent deux algorithmes d'apprentissage machine différents, soient la méthode de forêt aléatoire (*Random Forest, RF*) et le k-NN, afin de prédire la qualité de la nouvelle connexion optique. En testant leur modèle sur des données synthétiques, L. Barletta (March 2017) et Cristina Rottondi (2018) démontrent que le RF représente un bon compromis entre la précision de la prédiction et le temps d'exécution de l'algorithme.

Mata et al. (2017) développent un estimateur de qualité hybride utilisant l'algorithme SVM. Ainsi, les auteurs comparent dans un premier temps la longueur de la connexion optique avec les longueurs disponibles dans la base de données synthétique. Si celle-ci ne correspond à aucun critère prédéfini par l'expert, l'algorithme SVM est alors utilisé dans le but de prédire la qualité de la nouvelle connexion. Par la suite, la méthode SVM est comparée avec le RF et la régression logistique, en continuité de son étude (Mata et al., 2018a). Mata et al. (2018a) démontrent ainsi que l'algorithme RF représente un bon compromis en termes de temps

d'exécution et de précision de l'estimateur pour la prédiction de la qualité d'une nouvelle connexion optique.

Plusieurs études également comparent différents algorithmes afin de déterminer ceux présentant les meilleures performances pour l'estimation de la qualité d'une nouvelle liaison. Dans ce contexte, Aladin et Tremblay (2018) utilisent les mêmes attributs que ceux de la méthode de L. Barletta (March 2017) et reprennent les méthodes citées précédemment, à savoir le SVM, le RF et le K-NN dans le but d'estimer la classe de BER par rapport à un seuil de BER donné. En analysant la précision, le temps d'exécution ainsi que l'évolutivité des différents estimateurs implémentés et testés sur une base de données synthétique, Aladin et Tremblay (2018) démontrent que l'algorithme SVM surpasse les autres méthodes, bien qu'ayant un temps d'exécution plus important.

Les réseaux de neurones sont également utilisés pour l'estimation de la qualité d'une nouvelle connexion optique. En effet, T. Panayiotou utilise les réseaux de neurones afin de déterminer si la qualité de la nouvelle connexion optique à établir respecte le niveau requis défini par l'opérateur. Aussi, T. Panayiotou implémente-t-il dans ses premières études (2019a; 2017) un estimateur centralisé utilisant comme attributs les caractéristiques de la liaison afin de déterminer le niveau (acceptable ou non) du facteur Q de la nouvelle connexion optique. Ces attributs utilisés par T. Panayiotou sont le nombre d'amplificateurs, la longueur totale de la connexion optique, la longueur maximale des liaisons optiques, les informations sur le nœud de destination et la longueur d'onde de la nouvelle connexion optique à établir. Évalué à partir de données synthétiques, cet estimateur présente de meilleures performances, en terme de rapidité du temps d'exécution que des modèles d'estimateur de qualité traditionnels. Par la suite, T. Panayiotou implémente un réseau de neurones à une couche utilisant deux différents modèles (centralisés et distribués) basés sur le même principe que dans sa précédente étude. Son nouvel estimateur a pour but de choisir parmi les deux modèles celui qui estime le mieux la classe du BER de la nouvelle connexion à établir (2019b).

Sartzetakis, Christodoulopoulos et Varvarigos (2019) développent, quant à eux, une approche hybride. À partir d'une base de données synthétique qui prend en compte les effets non linéaires et linéaires de dégradation du signal ainsi que les paramètres des équipements, les auteurs génèrent premièrement de nouvelles données d'entrée en utilisant un modèle gaussien pour déterminer l'état de la connexion optique. Ces nouvelles données sont par la suite utilisées dans un modèle de régression non linéaire. Sartzetakis, Christodoulopoulos et Varvarigos (2019) parviennent donc, avec ce modèle, à estimer la qualité de transmission d'une nouvelle connexion optique en utilisant une base de données relativement petite (contenant 400 instances).

L'estimation de la qualité de transmission d'une nouvelle connexion optique peut également se faire par l'estimation du niveau de performance (plus précisément l'OSNR) de la nouvelle connexion à établir. C'est le cas de Proietti et al. (2019) qui développent un estimateur de qualité avec comme attributs d'entrée la puissance par canal ainsi que le niveau de bruit pour chaque liaison existante du réseau. Cet estimateur est implémenté à partir d'un réseau de neurones à deux couches et permet de déterminer l'OSNR de la nouvelle connexion à établir. Tout comme les précédentes méthodes, il est testé sur une base de données synthétique.

Dans le même contexte de prédiction du niveau de performance, Morais et Pedro (2018), implémentent premièrement des estimateurs basés sur différents algorithmes d'apprentissage machine (réseaux de neurones, K-NN, régression logistique et SVM) avec pour but de déterminer la qualité de la marge résiduelle de la connexion optique. Testés sur des données synthétiques, avec comme attributs les caractéristiques de la connexion optique, ils démontrent que l'estimateur utilisant les réseaux de neurones fournit de meilleures performances que les autres modèles. Ces attributs utilisés sont le nombre d'amplificateurs, la longueur totale de la connexion optique, l'information sur la source et la destination, la valeur moyenne de l'atténuation, le format de modulation, la valeur moyenne de la dispersion et la longueur maximale de la liaison. Par la suite, en utilisant les mêmes attributs, Morais et Pedro (2018)

utilisent l'estimateur implémenté à partir des réseaux de neurones pour estimer la valeur exacte de la marge résiduelle de la connexion optique.

Toujours dans l'estimation des valeurs exactes du niveau de performance des connexions optiques, Choudhury et al. (2018) font une étude comparative en utilisant l'algorithme RF et les modèles de régression linéaire afin de déterminer le BER d'une nouvelle connexion optique à établir. Les auteurs montrent que l'algorithme RF, entraîné à partir d'une base de données faible (2700 échantillons de BER), donne une meilleure estimation de la valeur du BER, comparativement au modèle de régression linéaire.

Vejdannik et Sadr (2020) utilisent les réseaux de neurones pour implémenter un modèle permettant d'estimer la puissance du bruit généré dans la liaison en utilisant comme attributs la longueur de la liaison, la puissance du canal et le format de modulation ainsi que la puissance et les formats de modulation des canaux voisins. À partir de ce modèle, ils calculent l'OSNR de la nouvelle connexion à établir.

De façon générale, les différentes méthodes utilisées ont été implémentées et testées sur des données synthétiques. Ils ont pour but d'estimer la qualité de transmission d'une nouvelle connexion optique en prédisant le niveau (bon ou mauvais) de qualité de la connexion optique déterminé à partir d'un seuil prédéfini par l'opérateur. En outre, d'autres méthodes ont pour but d'estimer la qualité de transmission dans un réseau à partir de la prédiction de la valeur du paramètre de performance (OSNR, puissance ou BER) de la nouvelle connexion optique. Cependant, ces approches utilisent une base de données très réduite. Ces contraintes sur les données utilisées (données synthétiques et données réduites) ne sont pas représentatives des conditions réelles dans le réseau. Par ailleurs, la plupart des attributs utilisés pour l'implémentation de ces estimateurs ont été basés sur les caractéristiques de la connexion optique. Il est bon de noter que hormis les paramètres internes du système de transmission et du réseau, tels que la puissance par canal, la longueur de la liaison, etc., d'autres facteurs externes (température, facteurs humains, etc.) peuvent également avoir un impact dans la

qualité de la transmission d'une connexion optique. Dans ce contexte, il devient important de développer des méthodes combinant à la fois les facteurs externes et internes pour l'estimation de la qualité de transmission d'une connexion optique. Par ailleurs, les approches implémentées à ce jour n'utilisent pas les données collectées sur le terrain sur de longues périodes d'observation. L'utilisation de celles-ci dans les modèles d'apprentissage machine pourrait permettre de tenir compte des effets saisonniers et des effets des interventions humaines dans le réseau afin de prédire les éventuelles dégradations de performance du réseau.

La prochaine section présente les principales méthodes utilisées pour détecter les anomalies survenant sur une connexion déjà établie.

1.4.2 Détection des pannes

Causées par des facteurs divers (vieillessement de la fibre optique, dégradation des équipements du réseau, effets environnementaux, etc.), les dégradations de performance des connexions optiques peuvent entraîner des pannes, ruptures de service et pertes de données. Les principales méthodes utilisées pour la gestion des pannes dans le réseau sont de type réactif. Celles-ci consistent à poser des actions de maintenance par les opérateurs lorsque des alarmes sont enclenchées à la suite d'une panne (Gu, Yang et Ji, 2020). Cependant, ces solutions réactives sont coûteuses pour les opérateurs, car ne prévenant pas la panne, elles ne permettent pas d'anticiper et donc d'éviter la perte de données dans le réseau. Les algorithmes d'apprentissage machine, qui utilisent les informations fournies par l'environnement du système ou les données collectées sur le terrain, pourraient permettre d'identifier les pannes avant que celles-ci ne surviennent, de poser des actions proactives, de diagnostiquer le problème présent sur le réseau. Ils représentent de ce fait des solutions permettant une gestion proactive des pannes. La gestion proactive des pannes peut être divisée en plusieurs catégories, à savoir : la prédiction de la panne, l'identification de la panne incluant sa classification (logicielle ou matérielle) et la gestion de la panne. Cette section présente les deux premiers groupes de gestion proactive des pannes (la prédiction et l'identification des pannes).

Ainsi, dans la catégorie de prédiction des pannes, Wang et al. (2017) utilisent l'algorithme SVM combiné avec le lissage exponentiel pour prédire les dégradations des équipements du réseau. Leur approche consiste dans un premier temps à déterminer les valeurs futures des attributs à mettre en entrée du modèle SVM en se basant sur les informations journalières fournies par l'opérateur. Ces attributs regroupent à la fois les paramètres physiques de l'équipement et les performances du réseau (consommation électrique, température interne et externe du circuit, puissance optique du laser, puissance optique, OSNR, etc.). Cette prédiction est effectuée avec la méthode de lissage exponentiel. Par la suite, l'algorithme SVM détecte la dégradation de l'équipement du réseau en se basant sur les données prédites par le lissage exponentiel.

Ruiz et al. (2016) quant à eux utilisent les réseaux bayésiens dans le but d'identifier les pannes dans une liaison. Dans leur étude, les pannes sont causées par deux effets : les interférences entre canaux et les pannes dues au filtrage étroit. Ainsi, à partir de données synthétiques de la puissance au récepteur de différentes liaisons optique d'un réseau, les caractéristiques temporelles de la puissance sont d'abord extraites (valeur moyenne, valeur maximale, valeur minimale, changement, tendance ...). À ces caractéristiques sont ensuite ajoutées les données synthétiques du pré-FEC BER. À partir de ces attributs, leur modèle détermine la probabilité d'avoir une panne sur la liaison et identifie la cause de cette panne.

Dans la continuité de l'étude précédente, Vela et al. (2017a; 2017b; 2017c) développent un modèle permettant de détecter et d'identifier les pannes survenant dans un réseau à partir de l'observation et de l'analyse du pré-FEC BER de chaque connexion optique. Ce modèle est implémenté en deux étapes. La première étape est l'étape de détection des pannes. Cette détection se fait à partir de l'implémentation d'un module appelé BANDO qui a pour but de détecter les anomalies et les variations soudaines du pré-FEC BER. Chacune de ces anomalies et variations est définie à partir d'un niveau de BER fixé par l'opérateur. La seconde étape est l'étape d'identification des pannes, implémentée à partir du module appelé LUCIDA utilisant

les réseaux bayésiens. Cette étape utilise comme attributs l'état fourni par le module BANDO, la puissance reçue, la valeur de la tendance du pré-FEC BER et la valeur de la périodicité du pré-FEC BER. Ainsi, à partir de données synthétiques, les auteurs ont pu implémenter des algorithmes permettant de générer une alerte lorsqu'une anomalie est détectée sur le pré-FEC BER.

Similairement, Shahkarami et al. (2018) développent une méthode permettant de détecter et d'identifier les pannes survenant sur un réseau. En effet, dans un premier temps, à partir de données de BER les caractéristiques statistiques de celles-ci (valeur maximale, valeur minimale, valeur moyenne et écart-type) sont extraites pour définir les attributs du modèle. Par la suite, le RF, l'algorithme SVM et un réseau de neurones sont utilisés pour détecter les anomalies sur le BER. Enfin, une fois la panne détectée, un réseau de neurones à deux couches est utilisé pour identifier la source de la panne.

Pour résumer, la plupart des méthodes d'apprentissage machine pour la détection des anomalies sont utilisées principalement pour détecter les défaillances logicielles observées dans le réseau. Celles-ci ne tiennent compte que des données synthétiques de BER et des paramètres du réseau. Tout comme les modèles d'estimation de la qualité de transmission, l'utilisation de données synthétiques ne représente pas les différents cas observés par l'opérateur. Le développement de modèles de prédiction de performance des connexions optiques et de détection d'anomalies entraînés avec des données réelles de terrain et intégrés dans les outils de gestion et de contrôle permettrait donc de déployer des solutions proactives de maintenance du réseau.

Dans les chapitres suivants, des modèles de classification de performance des connexions optiques sont proposés et des modèles prédictifs de performance basés sur des données de terrain sont proposés.

CHAPITRE 2

CLASSIFICATEUR DE PERFORMANCE D'UNE CONNEXION OPTIQUE

2.1 Introduction

L'utilisation des algorithmes d'apprentissage machine dans les estimateurs de qualité de transmission permet de réduire le temps de calcul et de prédire la performance d'un signal optique dans un contexte réel où tous les paramètres du système de transmission et de la connexion optique ne sont pas disponibles. Néanmoins, dans la littérature, les algorithmes d'apprentissage machine sont implémentés et testés soit en utilisant des données synthétiques (Musumeci et al., 2019b), soit en utilisant un petit échantillon de données réelles (Choudhury et al., 2018). Ils ne représentent ainsi pas tous les cas pouvant être observés dans les réseaux. Afin de pallier ce problème, ce chapitre présente donc deux classificateurs de performance de la connexion optique implémentés et testés à la fois sur des données réelles de performance optique et sur une plus grande base de données. Notons que ces classificateurs de performance sont implémentés en utilisant des techniques d'apprentissage machine existantes. Ces dernières sont optimisées et adaptées à notre contexte d'étude.

Ce chapitre est le résultat de travaux soumis dans un article de conférence, dans un article de journal et présentés à deux conférences :

- Allogba, S., et C. Tremblay. June 5-7, 2017. « Prediction of BER Trends in an Optical Link Using Machine Learning ». In 19th Photonics North Conf. (Ottawa, Canada, June 5-7, 2017).
- Allogba, S., et C. Tremblay. June 5-7, 2018. « Statistical analysis of performance monitoring data collected in an optical link ». In 20th Photonics North Conf. (Montreal, Canada, June 5-7, 2018).

- Allogba, S., et C. Tremblay. 2018. « K-Nearest Neighbors Classifier for Field Bit Error Rate Data ». In 2018 Asia Communications and Photonics Conference (ACP). (26-29 Oct. 2018), p. 1-3.
- Aladin, S., A. V. S. Tran, S. Allogba et C. Tremblay. 2020. « Quality of Transmission Estimation and Short-Term Performance Forecast of Lightpaths ». Journal of Lightwave Technology, vol. 38, no 10, p. 2807-2814.

Ce chapitre est divisé en deux parties. La première partie présente un estimateur de qualité de performance d'une connexion optique implémenté à partir de l'algorithme K-NN et testé sur des données réelles. En outre, il utilise également les caractéristiques externes de la liaison comme attributs supplémentaires afin de déterminer la qualité (« bonne » ou « mauvaise ») de la connexion optique (Allogba et Tremblay, 2018; June 5-7, 2018; Allogba et Tremblay, June 5-7, 2017). La deuxième partie présente l'impact des attributs sur la performance de l'estimateur de qualité (Aladin et al., 2020a).

2.2 Classificateur de performance utilisant l'algorithme K-NN

Cette section décrit le classificateur de BER basé sur l'algorithme K-NN. Notons que l'objectif est d'évaluer la possibilité d'utiliser un modèle simple d'apprentissage machine afin de caractériser les données de BER collectées dans une connexion optique et d'analyser l'impact des différents attributs sur la caractérisation du BER. Il serait possible d'utiliser d'autres techniques d'apprentissage machine plus sophistiquées et intéressant ultimement de faire une analyse comparative de performance avec notre modèle.

Le classificateur de BER proposé a été implémenté en suivant les trois étapes présentées dans la Figure 2.1.

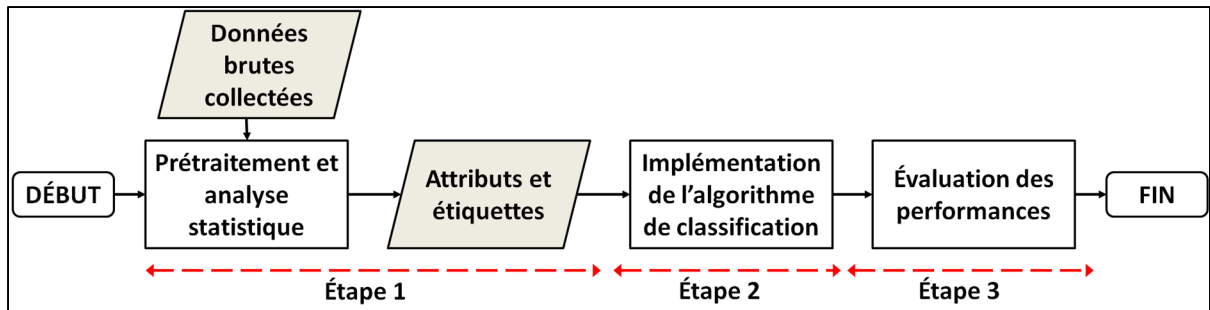


Figure 2.1 Étapes d'implémentation du classificateur de performance
Adaptée de Allogba et Tremblay (2018)

La première étape est le prétraitement et l'analyse statistique des données de BER collectées dans une connexion optique d'un réseau de production. Cette étape permet d'identifier et d'extraire les attributs et les étiquettes utilisés dans le classificateur. Par la suite, la seconde étape est la construction de l'algorithme K-NN pour classifier les données de BER. Finalement, la troisième étape est l'évaluation du classificateur à partir de l'analyse du temps de calcul et de la précision de prédiction.

2.2.1 Prétraitement des données

La liaison étudiée dans cette partie, présentée dans la Figure 2.2, est une connexion optique WDM de 230 km du réseau de CANARIE, avec 170 km en liaison enfouie et 60 km en liaison aérienne. La connexion optique WDM est une liaison qui permet de propager plusieurs signaux dans la même fibre optique. La connexion optique amplifiée transporte plusieurs canaux WDM. Elle comprend quatre segments (*spans*), trois amplificateurs optiques à l'erbium (*Erbium Doped Fiber Amplifier*, EDFA) et ROADM.

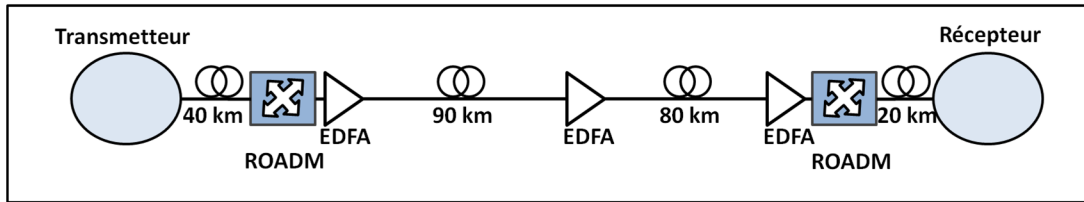


Figure 2.2 Connexion optique du réseau CANARIE
Adaptée de Allogba et Tremblay (2018)

Les données de performance, à savoir le pré-FEC BER ainsi que les paramètres de la connexion optique tels que la PMD et la PDL (*Polarization Dependent Loss*), ont été collectées toutes les secondes à partir d'un modem DP-QPSK (*Dual-Polarization Quadrature Phase Shift Keying*) à 40 Gbit/s. La base de données utilisée dans cette étude est constituée des données de pré-FEC BER. L'évolution temporelle du pré-FEC BER durant la période d'observation de 5 semaines, ainsi que la distribution des données, est présentée à la Figure 2.3. La base de données obtenue est donc composée de 1 790 413 données de pré-FEC BER étalées sur une période de 31 jours et variant entre $6,39 \times 10^{-9}$ et $1,47 \times 10^{-7}$, avec une valeur moyenne de $3,26 \times 10^{-8}$ et un écart-type de $1,25 \times 10^{-8}$.

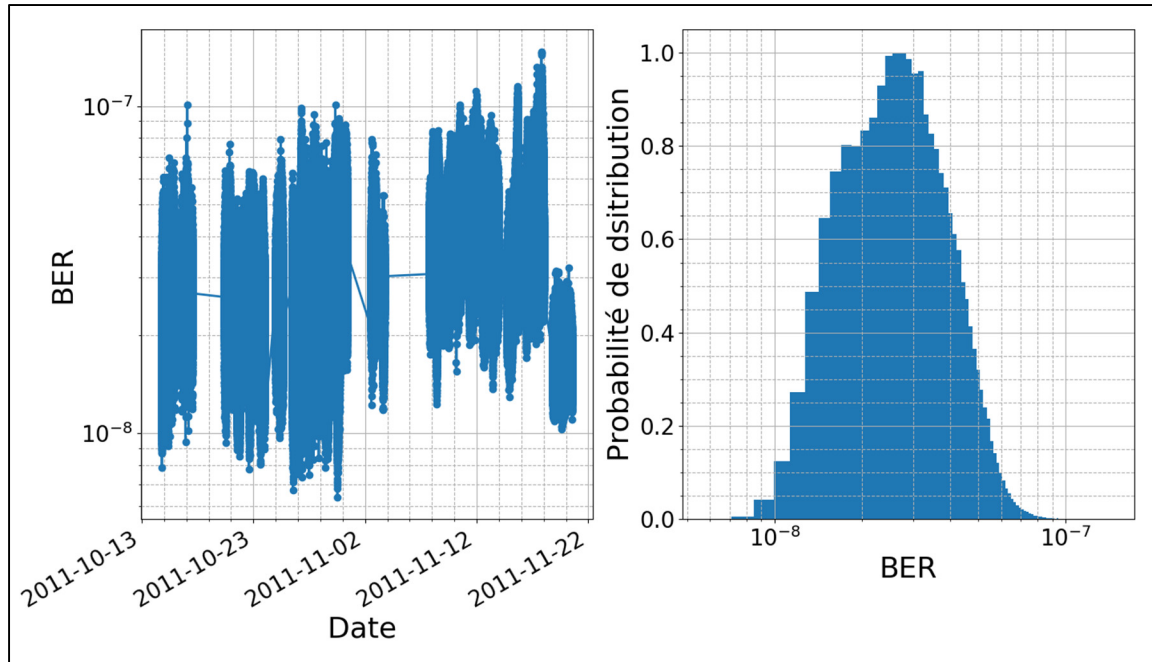


Figure 2.3 Évolution temporelle du BER ; Probabilité de distribution du BER

La phase de prétraitement des données est divisée en deux parties. La première partie consiste à définir les étiquettes à prédire par le modèle. À cet effet, dans un premier temps, les données brutes sont moyennées par minute dans le but de réduire le bruit dans les données brutes, formant une nouvelle base de données composée de 32 201 éléments de pré-FEC BER moyennés variant entre $1,09 \times 10^{-9}$ et $9,66 \times 10^{-8}$ avec un écart-type de $1,15 \times 10^{-8}$ la réduction du bruit. Par la suite, la base de données finale est construite en effectuant une moyenne mobile à l'heure sur les données moyennées à la minute. La fenêtre de la moyenne mobile est choisie arbitrairement.

En outre, le niveau de seuil qui permet de définir les étiquettes est $4,41 \times 10^{-8}$ et est calculé à partir de l'équation (2.1) :

$$BER_{th} = mean(BER) + \sigma(BER)_{min} \quad (2.1)$$

où :

- BER_{th} représente le seuil permettant de définir les étiquettes ;
- $mean(BER)$ représente la moyenne des données ;
- $\sigma(BER)_{min}$ représente l'écart-type des données moyennées par minute.

En pratique, le seuil peut être défini en se basant sur la marge du système fourni par les opérateurs. Les étiquettes sont donc déterminées en comparant la base de données finale avec le seuil. De cette comparaison, deux classes, représentant les étiquettes, sont identifiées, à savoir la classe 1 pour les « bons BER » et la classe 2 pour les « faibles BER ». La Figure 2.4 représente un exemple de détermination des étiquettes.

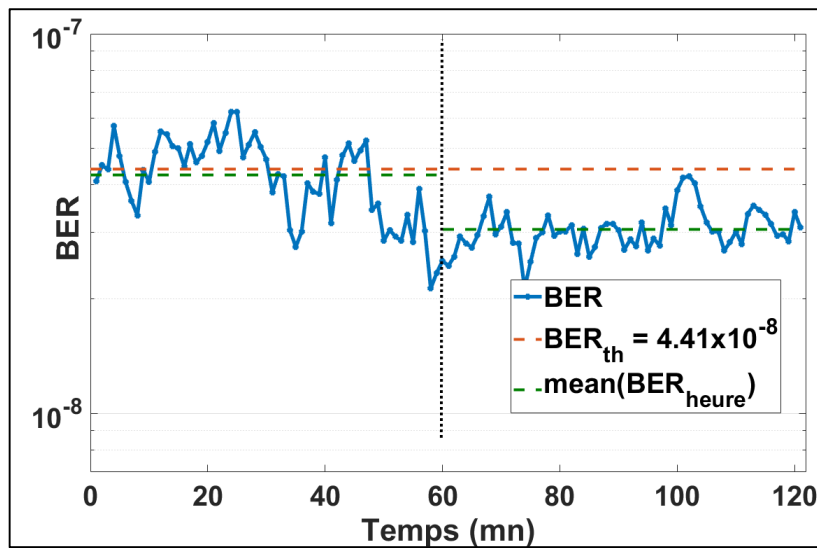


Figure 2.4 Exemple de représentation du seuil pour l'identification des étiquettes

Comme montrée dans la Figure 2.4, la classe 1 correspond à tous les éléments pour lesquelles la valeur moyenne prise sur une fenêtre d'une heure est inférieure au seuil BER_{th} prédéfini. De même, la classe 2 correspond à tous les éléments pour lesquelles la valeur moyenne prise sur une fenêtre d'une heure est supérieure au seuil BER_{th} prédéfini. La base de données finale contient donc 537 éléments, dont 342 éléments correspondants aux données de classe 1 et 142 éléments correspondants aux données de classe 2.

La deuxième partie du prétraitement des données est le lieu de déterminer les attributs utilisés dans le processus de classification. Afin d'augmenter les performances du classificateur, il est important que les attributs choisis aient un lien de dépendance avec les éléments de la base de données. Tremblay et al. (November 7-10, 2012); (July 3-6, 2017) présente une corrélation entre la température et les variations journalières de la PMD et de la PDL. Dans ce contexte, la température pourrait être considérée comme un attribut à prendre en compte par le modèle de classification. Ainsi, après avoir collecté la température durant la période d'observation (The Weather Company, 2014), une analyse de la température a été effectuée afin de déterminer un lien éventuel avec les données de BER. La Figure 2.5 présente les variations de la température et des données de BER de la base de données finale.

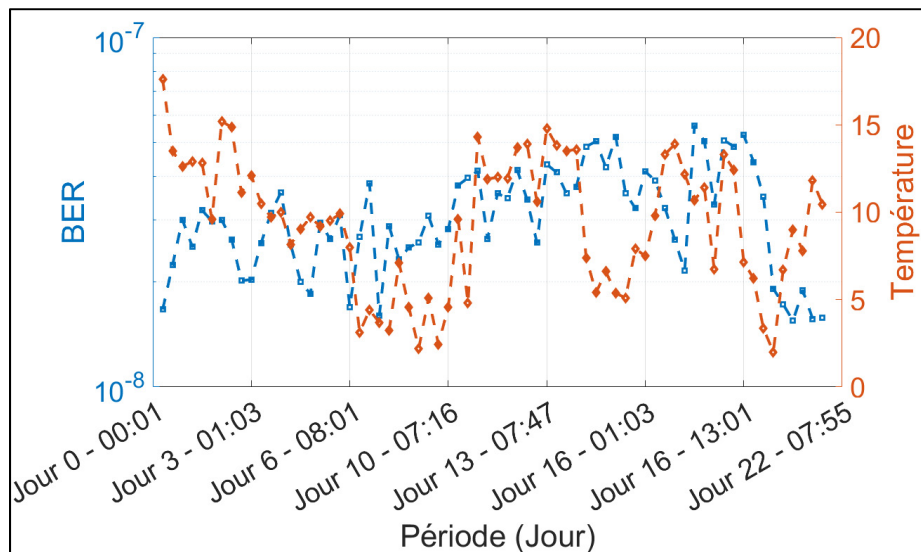


Figure 2.5 Évolution de la température du pré-FEC BER en fonction du temps durant une période d'observation de 22 jours
Adaptée de Allogba et Tremblay (June 5-7, 2018)

La température observée varie de -6°C à 18°C durant la période d'observation. Comme présentées dans la Figure 2.5, les variations des données de BER et de la température sont corrélées. En effet, pour des périodes données, lorsque la température est croissante, le BER est décroissant. De même, lorsque la température diminue le BER augmente. Afin de confirmer

cette analyse, une étude de corrélation a été effectuée en utilisant le test de Pearson. Ce test, réalisé sous MATLAB, a pour objectif de valider la corrélation entre les données horaires de BER et la température observée durant la période d'observation. Les résultats obtenus confirment donc la corrélation observée dans la Figure 2.5, avec une p-value de 0,0099 et un coefficient de corrélation de 75% (Allogba et Tremblay, June 5-7, 2018; June 5-7, 2017).

Les activités humaines pouvant également avoir un impact sur les performances du réseau, l'impact des différentes périodes de la journée, subdivisées en trois grands groupes à savoir le jour (de 08h00 à 15h59), le soir (de 16h00 à 23h59) et la nuit (de 00h00 à 07h59) a été également analysé. La Figure 2.6 représente la probabilité de distribution du BER pour ces trois périodes de la journée.

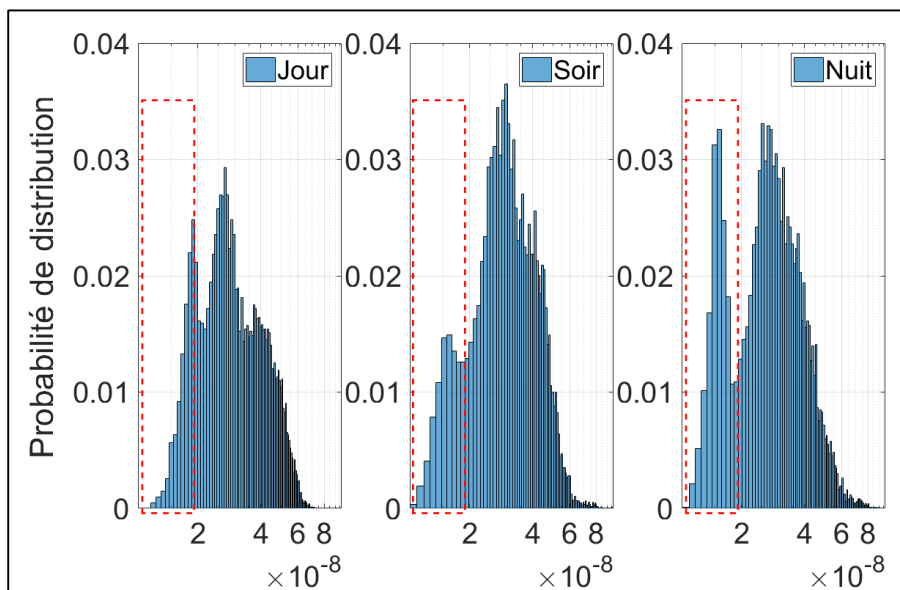


Figure 2.6 Probabilité de distribution du BER en fonction des périodes de la journée
Adaptée de Allogba et Tremblay (2018)

Les distributions sont différentes en fonction de la période de la journée considérée, principalement en observant les faibles valeurs de BER. En effet, les faibles valeurs de BER

sont observées plus durant la nuit que durant le jour ou le soir. Cette observation indique que la période de la journée peut décrire le comportement des variations du BER.

Les attributs à l'issue du prétraitement des données sont donc la température et les périodes de la journée.

2.2.2 Implémentation de l'algorithme de classification

L'objectif d'un classificateur est de prédire la classe d'une nouvelle observation. Dans le cadre de nos travaux, le classificateur implémenté a pour but de prédire la classe (« bon BER » ou « faible BER ») des données de BER. Il se base sur les données étiquetées, issues des données réelles de pré-FEC BER collectées dans une connexion optique, et des attributs déterminés à partir de l'étape de prétraitement des données. Par ailleurs, il peut être utilisé dans les outils de détection des anomalies afin de déceler les dégradations survenant sur la liaison. Cette partie décrit l'algorithme utilisé pour le modèle de classification ainsi que les étapes de son implémentation.

L'algorithme utilisé pour l'implémentation du classificateur est l'algorithme K-NN. Bien que moins performant que les autres algorithmes de classification, l'algorithme K-NN est plus rapide, car n'ayant pas besoin de généraliser sa méthode avec les données de la base d'apprentissage (Morais et Pedro, 2018; Singh, Thakur et Sharma, 2016).

Comme décrit dans la partie 1.3.2.2, l'algorithme K-NN est un algorithme non paramétrique utilisé pour déterminer l'étiquette d'une nouvelle donnée en se basant sur celles de la base de données ayant déjà des étiquettes. Le choix de l'étiquette de la nouvelle donnée est effectué par vote majoritaire sur les étiquettes de ses K données voisines (ou de ses plus proches éléments). Les données voisines sont déterminées, dans le cadre de cette étude, en calculant la distance euclidienne entre chacun des éléments. La Figure 2.7 présente un exemple de fonctionnement de l'algorithme K-NN.

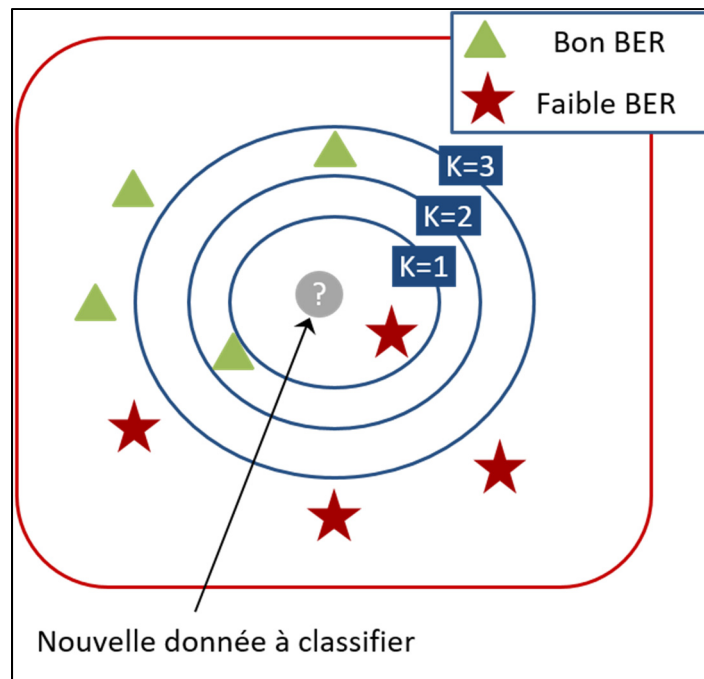


Figure 2.7 Principe de fonctionnement de l'algorithme K-NN

La nouvelle donnée à classifier X dépend de l'étiquette associée à chacun de ses éléments voisins. Ainsi, dans la Figure 2.7, la classe de X peut-elle être « faible BER », « bon BER » ou « faible BER » et « bon BER » dans les cas respectivement où $K=1$, $K=2$ et $K=3$. En effet, dans le cas $K=1$, le seul élément proche de la donnée X a comme étiquette « faible BER ». De même, dans le cas $K=3$, deux éléments sur les trois éléments proches de X sont de classe « faible BER ». Par contre, dans le cas $K=2$, les deux éléments proches de X sont de classes différentes. La classe de X peut donc être l'une des deux classes. Ce genre de cas peut se produire généralement pour des valeurs de K paires. Idéalement, il conviendrait de choisir des valeurs impaires pour le nombre de voisins afin d'éviter de se retrouver dans des situations incertaines.

Comme présentée dans la Figure 2.7, la classe d'une nouvelle donnée non étiquetée est fonction de la classe prédominante de ses voisins. Le choix du paramètre K a donc un impact sur les performances de l'algorithme K-NN. De ce fait, il convient d'optimiser le nombre de

voisins K à utiliser pour déterminer les nouvelles étiquettes. La détermination du paramètre K peut se faire soit en utilisant la méthode de validation croisée, soit en comparant le taux de précision trouvée à partir de différentes valeurs de K dans la base d'apprentissage. Dans notre cas, le paramètre K a été déterminé en utilisant la méthode de validation croisée à 10 plis. Celle-ci consiste à faire varier l'algorithme K -NN pour différentes valeurs de K variant de 1 à 20 et choisir la valeur de K donnant le plus fort taux de classification. La base de données est divisée en 10 échantillons de 54 éléments. Le modèle apprend à partir des 9 échantillons non sélectionnés, soit 483 éléments de la base de données. Par la suite, le paramètre K est testé sur le dixième échantillon et le taux d'erreur de classification est calculé. Cette étape est répétée jusqu'à ce que les 10 échantillons soient utilisés comme base de test. Le paramètre K choisi est déterminé à partir de la moyenne du taux d'erreur de classification issue de chaque essai. La Figure 2.8 présente la variation du taux de précision en fonction du nombre de voisins choisis K paramétré dans le classificateur.

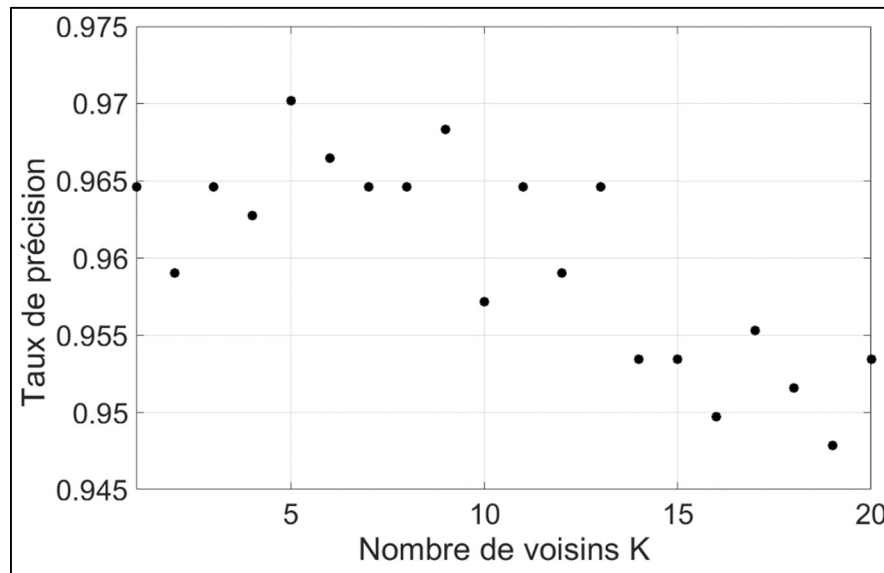


Figure 2.8 Variation du taux de précision en fonction du nombre de voisins K

Adaptée de Allogba et Tremblay (2018)

La meilleure valeur de K est obtenue pour K = 5 avec un taux de précision de 97%. Par ailleurs, plus le nombre de voisins augmente, plus la précision diminue.

Une fois le paramètre K défini, la validation du modèle de classification a été effectuée en utilisant 3 scénarios différents basés sur le choix des attributs :

- scénario 1 : température et période de la journée,
- scénario 2 : valeurs moyennes et maximales des données horaires de BER,
- scénario 3 : combinaison des attributs du scénario 1 et 2, à savoir la température, la période de la journée, les valeurs moyennes et maximales des données horaires de BER.

Chacun de ces modèles est testé en utilisant soit la validation croisée à 10 plis, soit la méthode de répartition de la base de données avec le ratio 80/20. Notons également que la base de test pour la validation des modèles correspond à 20% de la base de données. Le Tableau 2.1 représente les performances obtenues pour chacun des modèles implémentés.

Tableau 2.1 Évaluation des performances du classificateur de BER
Adapté de Allogba et Tremblay (2018)

Méthode	Métrique	Paramétrage des attributs		
		Scénario 1	Scénario 2	Scénario 3
Répartition 80/20	Taux de précision (%)	55,56	87,04	96,3
	Faux positifs (%)	12,04	10,44	1,54
	Temps de calcul (s) ¹	0,013	0,016	0,016
Validation croisée à 10 plis	Taux de précision (%)	73	90,7	97,8
	Faux positifs (%)	14,94	7,01	0,77
	Temps de calcul (s) ¹	0,74	0,76	0,76

¹Processeur Single AMD Phenom™ II X4 B95 avec 7,5 Gb mémoire RAM

Trois métriques sont observées dans l'évaluation des performances, comme présentées dans le Tableau 2.1, à savoir : le taux de précision pour indiquer le nombre de prédictions exactes faite

par le modèle; le taux de faux positif pour décrire le nombre de fois où le modèle a prédit une donnée de classe « faible BER » alors que la classe réelle est « bon BER » et le temps de calcul pour déterminer la vitesse d'exécution des modèles. Le taux de précision et le taux de faux positif sont déterminés en utilisant la matrice de confusion présentée à la Figure 2.9.

Étiquette réelle	2	86	56	24	116	9	139
	1	336	59	371	26	386	3
		1	2	1	2	1	2
		Étiquette prédite					

Figure 2.9 Matrice de confusion du classificateur de BER utilisant la méthode par validation croisée à 10 plis

Adaptée de Allogba et Tremblay (2018)

La meilleure classification est donc obtenue en utilisant tous les attributs (température, période de la journée, valeurs moyennes et valeurs maximales). En effet, le taux de prédiction est de 96,3% pour la méthode utilisant la répartition des données suivant le ratio 80/20 et de 97,8% pour la méthode utilisant la validation croisée à 10 plis. Quant au taux de faux positif, il est de 1,54% et 0,77%, respectivement pour la méthode utilisant la répartition des données suivant le ratio 80/20 et pour la méthode utilisant la validation croisée à 10 plis. De plus, en analysant le Tableau 2.1, toutes les performances sont meilleures en utilisant la méthode de validation croisée à 10 plis. Cependant, cette méthode est plus lente que la méthode de répartition des données avec le ratio 80/20, avec une vitesse pouvant aller jusqu'à 0,76 s contre 0,016 s (observée dans le scénario 3).

Par ailleurs, plus le nombre d'attributs est élevé, plus le modèle est plus performant et plus le temps de calcul est élevé. Ainsi, en faisant un compromis entre le taux de précision et le temps de calcul, la méthode utilisant la répartition des données selon le ratio 80/20 (avec une base de données de 537 instances), avec un taux de précision de 96,3% et un rapide temps d'exécution

(0,016s), présente un potentiel pour des scénarios réels d'implémentation d'outils de contrôle proactifs de la dégradation des performances à travers une classification.

Afin d'améliorer les performances des classificateurs de BER utilisant les algorithmes d'apprentissage machine, des travaux portant sur l'augmentation de la taille de la base de données, la variation des attributs ou l'utilisation de différentes techniques d'apprentissage peuvent être appliqués. Dans ce contexte, la prochaine partie présente nos travaux portant sur l'étude des attributs sur la performance du classificateur de BER.

2.3 Étude de l'impact des attributs dans les classificateurs

Dans cette section, le classificateur implémenté est un estimateur de qualité de liaison (bonne ou mauvaise qualité). Il permet ainsi de déterminer la possibilité d'établir une nouvelle connexion en utilisant les informations du réseau telles que la longueur du lien, la longueur en chaque amplificateur, les pertes, etc. Cependant, dans un contexte réel, ces informations de la liaison peuvent être incomplètes ou inexactes à cause des erreurs possibles lors de la collecte de données. Ainsi, l'utilisation des méthodes d'apprentissage machine peuvent permettre de réduire le nombre d'attributs à utiliser afin de minimiser l'impact du manque de données dans les performances de l'estimateur de qualité. Pour ce faire, cette section est divisée en deux parties. La première partie présente les modèles utilisés pour estimer la qualité de la liaison. La seconde partie présente la base de données utilisée dans l'estimateur de qualité, la méthode de réduction des attributs et l'impact de ceux-ci sur la performance de l'estimateur.

2.3.1 Présentation de l'estimateur de qualité

Deux modèles, basés sur deux méthodes d'apprentissage machine, à savoir le SVM et les réseaux de neurones, ont été implémentés pour estimer la qualité de la liaison (Aladin et al., 2020a). Contrairement à l'algorithme K-NN qui est sensible à la répartition de la base de données utilisée avec la notion de voisinage, l'algorithme SVM et les réseaux de neurones sont

peu sensibles à la base de données. En effet, une fois les modèles construits, l'ajout d'une nouvelle instance dans la base de données n'a aucun impact sur les performances du modèle.

L'algorithme SVM implémenté dans le premier modèle est un algorithme de classification supervisé principalement utilisé pour les problèmes avec deux classes. L'algorithme SVM a pour but de distinguer les deux classes de la base de données à l'aide d'un hyperplan qui agira comme une frontière entre les deux classes. Cet hyperplan est construit de telle sorte à ce que la distance entre les deux classes soit maximale. Cette distance est également appelée marge de séparation. L'emplacement de la marge de séparation est déterminé en utilisant un sous-ensemble d'observations de la base d'apprentissage, appelé vecteurs de support.

Tout comme l'algorithme K-NN, le SVM est également construit à partir de différents paramètres tels que le noyau, l'équilibrage C permettant de contrôler l'influence des vecteurs supports et de minimiser les erreurs de classification ainsi que γ définissant le noyau utilisé. Dans le cadre de cette étude, le noyau gaussien est celui choisi pour l'implémentation de l'estimateur. Par ailleurs, afin d'optimiser les autres paramètres, la base de données a été divisée selon le ratio 80/20 puis la méthode de validation croisée à 5 plis a été utilisée dans la base d'apprentissage. Ainsi, les paramètres trouvés sont 1×10^3 et 0,5102, respectivement pour le paramètre C et le paramètre γ .

Quant au deuxième modèle proposé utilisant les réseaux de neurones (*Artificial Neural Networks, ANN*), il se base sur le vecteur d'entrée X (attributs) pour calculer la valeur de sortie $f(X)$, correspondant aux étiquettes associées, à partir d'un ensemble de poids w et de biais b . Par ailleurs, le réseau de neurones est composé de plusieurs couches qui peuvent augmenter le taux de prédiction du modèle à travers sa capacité à modéliser la base de données. Afin d'ajuster le poids et le biais utilisé par le réseau de neurones, la base de données est également divisée en suivant le ratio 80/20. La base d'apprentissage est utilisée pour sélectionner les hyper-paramètres optimaux permettant de minimiser le taux d'erreur, tel que le pas d'apprentissage, la taille des couches de neurone et le nombre de couches cachées. Ainsi, la

détermination des hyper-paramètres dans la phase d'apprentissage a permis de fixer le nombre de couches cachées à 2, la taille des couches cachées à 512 neurones pour la première couche cachée et 256 neurones pour la deuxième couche cachée et le pas d'apprentissage à 0,00074. Par ailleurs, le modèle ANN a été implémenté en utilisant l'optimiseur Adam pour des mini lots de 512 échantillons. De plus la fonction *Leaky ReLu* a été utilisée comme fonction d'activation des neurones à l'exception de la dernière couche de neurones qui a été conçue avec la fonction sigmoïde.

2.3.2 Étude de l'impact des attributs

La base de données utilisée pour l'estimation de qualité de la liaison est une base synthétique construite à partir d'un outil de générateur de données basé sur le modèle de bruit gaussien Aladin et Tremblay (2018). Cet outil permet d'estimer le BER en fonction du bruit linéaire et non linéaire ainsi que des caractéristiques de la liaison présentées dans le Tableau 2.2. Il a donc permis d'avoir une base de données synthétique de 38 400 éléments.

À partir de cette base de données synthétique, l'analyse de l'impact des attributs sur l'estimateur de qualité a été effectuée dans le but de refléter la disponibilité des paramètres dans un scénario réel.

Tableau 2.2 Paramètre de l'outil de générateur de données

	Paramètres	Valeurs
Paramètres variables	Longueur du lien	[80 ; 7500 km]
	Longueur des amplificateurs	80 ; 100 ; 120 ; 150 km

	Nombre d'amplificateurs	[1 ; 50]
	Format de modulation	PM-BPSK ; PM-QPSK ; PM-16QAM ; PM-64QAM
	Puissance	[-10 ; 5 dBm]
	Débit	40 ; 50 ; 100 Gbit/s
Paramètres fixes	Atténuation de la fibre	0,2 dB/km
	Coefficient de dispersion	-21 (ps) ² /km
	Coefficient non linéaire	1,3 W ⁻¹ km ⁻¹
	Bruit de l'amplificateur	5 dB
	Fréquence centrale	193,1 THz
	Bruit	32 GHz
	Bande passante de référence	12,5 GHz
	Nombre de canaux	9
	Débit en bauds	32 Gbaud

Pour ce faire, dans un premier temps, le taux de précision de l'estimateur SVM a été évalué en considérant chaque attribut séparément. Ces tests ont été réalisés sous MATLAB. Les attributs considérés sont les paramètres variables utilisés dans l'outil de générateurs de données. Ainsi, 6 modèles ont été définis à partir de ces attributs. Ces modèles ont été optimisés et testés en utilisant la méthode de validation croisée à 10 plis. Plus l'attribut a un impact sur la qualité de la liaison à estimer, plus le taux de précision est élevé. La Figure 2.10 présente les résultats du taux de précision pour chacun des attributs.

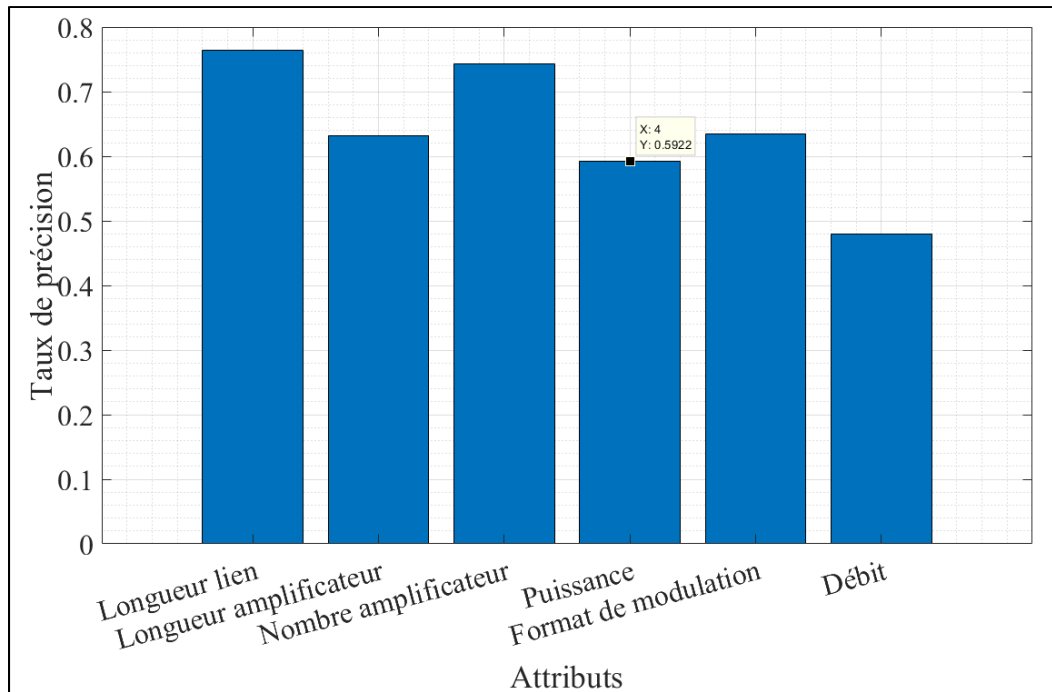


Figure 2.10 Taux de précision en fonction du choix des attributs
Adaptée de Aladin et al. (2020a)

Considérant la Figure 2.10, l'attribut ayant le plus d'impact sur la qualité de la liaison est la longueur de la liaison, avec un taux de précision de 76,47%, suivie du nombre d'amplificateurs, du format de modulation, de la longueur des amplificateurs, de la puissance et du débit.

Par la suite, afin de réduire le nombre d'attributs, une analyse de corrélation, réalisée sur MATLAB et utilisant le test de corrélation de Pearson a été effectué entre les différents attributs. Cette analyse a pour but de ne garder qu'un seul attribut parmi les attributs les plus fortement corrélés. Le test de Pearson a révélé un coefficient de corrélation de 0,33 entre la longueur des amplificateurs et le nombre des amplificateurs et de 0,93 entre la longueur des amplificateurs et la longueur de la liaison. Par ailleurs, dans le cadre de cette étude, la longueur entre chaque amplitude est considérée comme égale. Par conséquent, les 4 paramètres résultants des analyses d'impact des classificateurs sont la longueur de la liaison, le nombre d'amplificateurs, la puissance et le format de modulation. Dans un contexte réel où la longueur

entre les amplificateurs est variable (sauf dans les liaisons sous-marines), le choix des 4 paramètres pourrait être différent, car la valeur du SNR d'une liaison dépend de la plus grande longueur d'amplificateur dans un lien.

Pour évaluer l'impact du choix des attributs sur les deux modèles d'estimateur de qualité construit, trois scénarios ont été représentés :

- le premier comprenant tous les attributs (longueur de la liaison, longueur des amplificateurs, nombre des amplificateurs, format de modulation, puissance et débit),
- le deuxième comprenant les 4 meilleurs attributs (longueur de la liaison, nombre d'amplificateurs, format de modulation et puissance),
- le dernier comprenant les trois meilleurs attributs (longueur de la liaison, nombre d'amplificateurs, format de modulation).

Par ailleurs, afin de tenir compte du fait que le nombre d'éléments de la classe « mauvaise qualité » est plus élevé (76,67% de la base de données) que le nombre d'éléments de la classe « bonne qualité », une attribution de coûts a été effectuée afin de rééquilibrer les erreurs de classification. Cette attribution de coûts est basée sur une matrice de coûts insérée dans le paramètre *cost* fourni par la librairie *fitcsvm* de MATLAB.

Les métriques observées sont le taux de précision et le temps de calcul selon le Tableau 2.3. Il est important de noter que le temps de calcul ne tient pas compte du temps d'apprentissage du modèle. Il correspond à la vitesse de prédiction du modèle.

Comme présentés dans le Tableau 2.3, les estimateurs donnent de meilleures performances lorsque tous les attributs sont pris en compte avec 99,56% pour le modèle utilisant la méthode ANN et 99,38% pour le modèle SVM.

Tableau 2.3 Évaluation de performance de l'estimateur de qualité

Modèle	Métrique	Paramétrage des attributs
--------	----------	---------------------------

		Scénario 1	Scénario 2	Scénario 3
ANN	Taux de précision (%)	99,56	92,30	85,03
	Temps de calcul (ms) ¹	0,276	0,214	0,268
SVM	Taux de précision (%)	99,38	93,30	88,52
	Temps de calcul (ms) ¹	3,45	1,01	0,935

¹Processeur Intel® Core™ i5-8600K 3.6 GHz CPU, 16 GB RAM et un GTX 970 GPU.

En outre, plus le nombre d'attributs diminue, plus le taux de précision diminue. En effet, pour le scénario 2, le taux de précision est de 92,30% et 93,30%, respectivement pour le modèle ANN et le modèle SVM. Quant au scénario 3 correspondant à l'utilisation des 3 meilleurs attributs, le taux de précision est de 85,02% pour le modèle ANN et 88,52% pour le modèle SVM.

Cependant, le modèle SVM est plus performant que le modèle ANN lorsque le nombre d'attributs diminue. Plus le nombre d'attributs est élevé, plus le temps de calcul est important, dans le cadre du modèle SVM.

2.4 Conclusion

Ces travaux ont permis d'évaluer la possibilité d'implémenter les modèles d'apprentissage machine pour classifier ou estimer la qualité d'une connexion optique existante ou d'une nouvelle connexion à établir. Premièrement, avec la méthode K-NN, il a été question d'implémenter un modèle permettant de classifier le niveau du BER en utilisant des données réelles d'une connexion optique existante (valeurs maximales et moyennes) ainsi que les facteurs externes à la liaison, à savoir la température, les activités humaines définies par la

période de la journée. Bien qu'ayant des performances moins bonnes que des algorithmes d'apprentissage machine plus puissants tels que les réseaux de neurones ou le SVM, ce modèle a permis de classifier les différentes données avec un taux de précision allant jusqu'à 97,8%.

Par la suite, la deuxième étude, quant à elle, a permis d'implémenter un estimateur de qualité de liaison. Cet estimateur a été construit en utilisant une base de données synthétique, mais plus grande que l'étude précédente (38 400 éléments). De plus, les attributs utilisés dans cette partie sont les caractéristiques d'une connexion optique, contrairement à l'étude précédente qui était basée sur les paramètres externes de la liaison. En outre, une étude a été faite sur les attributs afin de se placer dans un contexte réel dans lequel toutes les données pouvaient ne pas être disponibles. Les deux modèles utilisés pour cet estimateur ont donc montré qu'il était possible d'estimer la qualité de la connexion optique dans le cas où le nombre d'attributs était réduit, avec un taux de précision allant de 93.30% (cas avec 4 attributs) à 88.52% (cas avec 3 attributs).

Ces deux études peuvent être utilisées en remplacement des modèles analytiques d'estimation de la qualité d'une connexion optique.

Il est toutefois important de noter que, lors de l'évaluation des performances des modèles, aucune analyse statistique n'a été effectuée sur les métriques. Une telle analyse permettrait éventuellement d'évaluer la robustesse des métriques trouvées pour chacun des modèles implémentés.

Toujours dans l'optique d'implémenter des outils de contrôle proactif des réseaux, les prochains chapitres portent sur la prédiction des valeurs de performance des connexions optiques.

CHAPITRE 3

OUTILS DE PRÉDICTION DES PERFORMANCES D'UNE CONNEXION OPTIQUE

3.1 Introduction

Les algorithmes d'apprentissage peuvent être également utilisés pour la prédiction des métriques de performance du réseau, telles que le BER, le facteur de qualité, le trafic, etc. Francesco Musumeci et al. (2018) présente une revue de l'état de l'art sur l'utilisation des algorithmes d'apprentissage machine dans diverses applications, notamment la prédiction du trafic du réseau. Dans ce contexte, Balanici et Pachnicke (2019) utilisent dans leurs travaux une variante des réseaux de neurones (*Nonlinear AutoRegressive Neural Network, NARNN*) sur des données synthétiques pour prédire le trafic dans le réseau, équivalent au taux d'utilisation du réseau. De même, Fernández et al. (2013) utilise la méthode de régression ARIMA (*AutoRegressive Moving Average*) pour déterminer la matrice de trafic dans l'intervalle de temps (*time slot*) suivant. Mo et al. (2018a) quant à lui propose le LSTM (*Long Short-Term Memory*), une autre variante des réseaux de neurones, pour prédire le trafic requis pour le ROADM 30 minutes à l'avance. Cependant, ces méthodes ne sont pas représentatives des situations réelles. En effet, la base de données utilisée est soit générée à partir de données synthétiques, soit générée sur un petit échantillon de données réelles. De plus, la prédiction effectuée a été faite sur des horizons courts (correspondant généralement à l'instant suivant).

Dans ce contexte, ce chapitre porte sur l'implémentation de modèles de prédiction de performance sur des horizons allant jusqu'à 4 jours. Ces modèles de prédiction sont implémentés en utilisant les algorithmes LSTM et GRU (*Gated Recurrent Unit*). Le choix de ces algorithmes est justifié par le fait que ceux-ci sont plus robustes aux effets de non-stationnarité et de non saisonnalité des données. La méthode d'optimisation et d'implémentation des modèles est présentée à la Figure 3.1. Les modèles sont basés sur des

algorithmes d'apprentissage machine construits et testés à partir des données réelles collectées dans des liaisons optiques d'un réseau de production. Les travaux présentés dans ce chapitre ont fait l'objet de deux articles de conférence et un article de revue :

- C. Tremblay, S. Allogba and S. Aladin, "Quality of transmission estimation and performance prediction of lightpaths using machine learning," 45th European Conference on Optical Communication (ECOC 2019), Dublin, Ireland, 2019, pp. 1-3, doi: 10.1049/cp.2019.0757.
- Aladin, S. Allogba, A. Tran, and C. Tremblay, "Recurrent Neural Networks for Short-Term Forecast of Lightpath Performance," in Optical Fiber Communication Conference (OFC) 2020, OSA Technical Digest (Optical Society of America, 2020), paper W2A.24.
- S. Aladin, A. V. S. Tran, S. Allogba and C. Tremblay, "Quality of Transmission Estimation and Short-Term Performance Forecast of Lightpaths," in Journal of Lightwave Technology, vol. 38, no. 10, pp. 2807-2814, 15 May 2020, doi: 10.1109/JLT.2020.2975179.

Ce chapitre est subdivisé en quatre parties. La première partie présente l'objectif de ce chapitre ainsi que les étapes d'implémentation des modèles de prédiction présentés à la Figure 3.1. La deuxième partie présente les données utilisées pour la construction et l'évaluation des modèles. Elle décrit également les méthodes de prétraitement des données. La troisième partie présente les modèles de prédiction de performance construits à partir de deux variantes des réseaux de neurones : le LSTM et le GRU) (Aladin et al., 2020a; Aladin et al., 2020b; Tremblay, Allogba et Aladin, 2019). Finalement, la quatrième partie du chapitre inclut une analyse comparative des performances des différents modèles de prédiction.

3.2 Objectif

L'objectif de ce chapitre est d'évaluer la possibilité de prédire le SNR d'une connexion optique pour des horizons allant jusqu'à 4 jours à partir des données historiques de SNR de la connexion. Deux bases de données issues de deux connexions optiques différentes ont donc

été utilisées. Les modèles utilisés pour la prédiction des données de SNR sont issus de deux variantes de réseaux de neurones. La Figure 3.1 présente les étapes d'implémentation des modèles de prédiction.

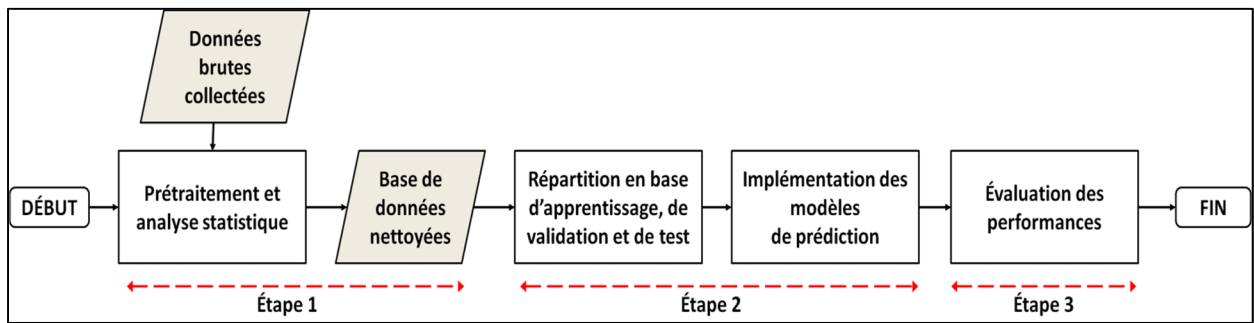


Figure 3.1 Procédure d'implémentation des modèles prédictifs

Chacune de ces étapes est décrite en détail dans les sections suivantes :

- Prétraitement des données : cette étape correspond à une phase de nettoyage des données. Rappelons que deux bases de données différentes sont utilisées pour l'évaluation de la prédiction des données de SNR. L'étape de prétraitement des données inclut donc la présentation des deux bases de données utilisées pour l'entraînement des modèles, le remplacement des données manquantes, l'analyse statistique et la génération de la base des données. Cette étape est présentée à la section 3.3.
- Implémentation des modèles : les modèles utilisés pour la prédiction des données de SNR sont le LSTM et le GRU. L'implémentation de ces modèles est présentée à la section 3.4 et comprend la description générale des algorithmes LSTM et GRU et la phase d'optimisation des hyper-paramètres.
- Évaluation des performances : elle est présentée dans la section 3.5. Par ailleurs, afin de valider la performance de nos modèles de prédiction, trois scénarios ont été testés. Cette étape présente donc les performances des modèles pour chacun des scénarios.

3.3 Génération de la base de données

3.3.1 Caractéristiques des deux connexions optiques du réseau de CANARIE

Les modèles de prédiction de performance présentés dans ce chapitre ont été développés à partir des données de deux connexions optiques (appelées A et B) du réseau de CANARIE. Les bases de données correspondantes A et B ont été utilisées pour le développement et la validation des modèles de prédiction du SNR sur des horizons de quelques jours.

Tableau 3.1 Caractéristiques des deux bases de données de terrain de deux connexions optiques utilisées pour les modèles de prédiction de performance

	CANARIE	
	Connexion optique A Base de données A	Connexion optique B Base de données B
Données extraites	BER	BER
Taille de la base de données	38 203	53 300
Données manquantes	1 020	7 150
Période d'observation	13 mois	18 mois

Le Tableau 3.1 récapitule les informations sur les deux bases de données. Les données sont issues de deux connexions optiques déployées sur deux liaisons optiques de 1300 km (connexion A et connexion B) du réseau de CANARIE. Chacune de ces liaisons est composée de 13 sites d'amplification espacés de 96 km environ et d'un site de ROADM à commutateurs sélectifs en longueur d'onde (WSS). En outre, plusieurs métriques de performance des deux

canaux PM-QPSK à 100 Gbit/s sont collectées par le système de contrôle du réseau à une fréquence d'échantillonnage de 15 minutes. À partir de ces métriques de performance, des données de BER ont été récupérées durant deux périodes d'observation différentes. Ces données de BER sont par la suite utilisées pour construire les bases de données A et B, composées des données de BER converties par la suite en données de SNR à partir de l'équation (3.1) :

$$SNR = 10 \log \left((\sqrt{2} * ierfc(2 * BER))^2 \right) \quad (3.1)$$

où :

- SNR s'exprime en dB ;
- BER représente les données de BER collectées toutes les 15 min ;
- ierfc représente l'inverse de la fonction d'erreur complémentaire.

L'évolution temporelle du SNR de la connexion optique A (base de données A) durant la période d'observation de 13 mois est illustrée à la Figure 3.2. Les valeurs de SNR sont comprises entre 7,96 dB et 10,35 dB avec une valeur moyenne de 9,91 dB et un écart-type de 0,30 dB.

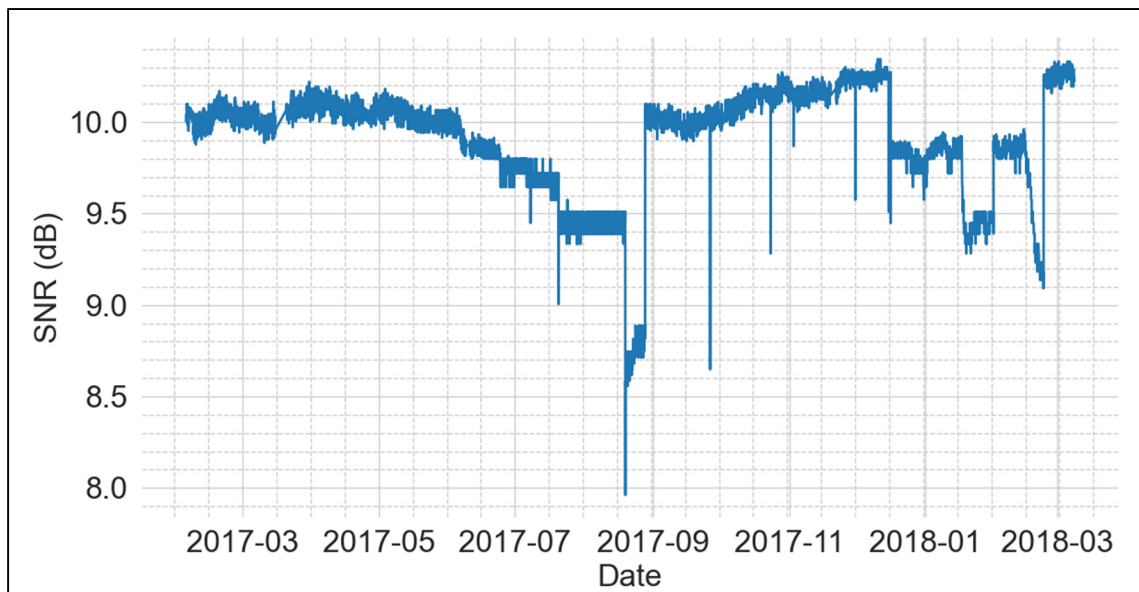


Figure 3.2 Évolution temporelle du SNR de la connexion optique A (base de données A) durant la période d'observation de 13 mois
Tirée de Aladin et al. (2020a); Tremblay, Allogba et Aladin (2019)

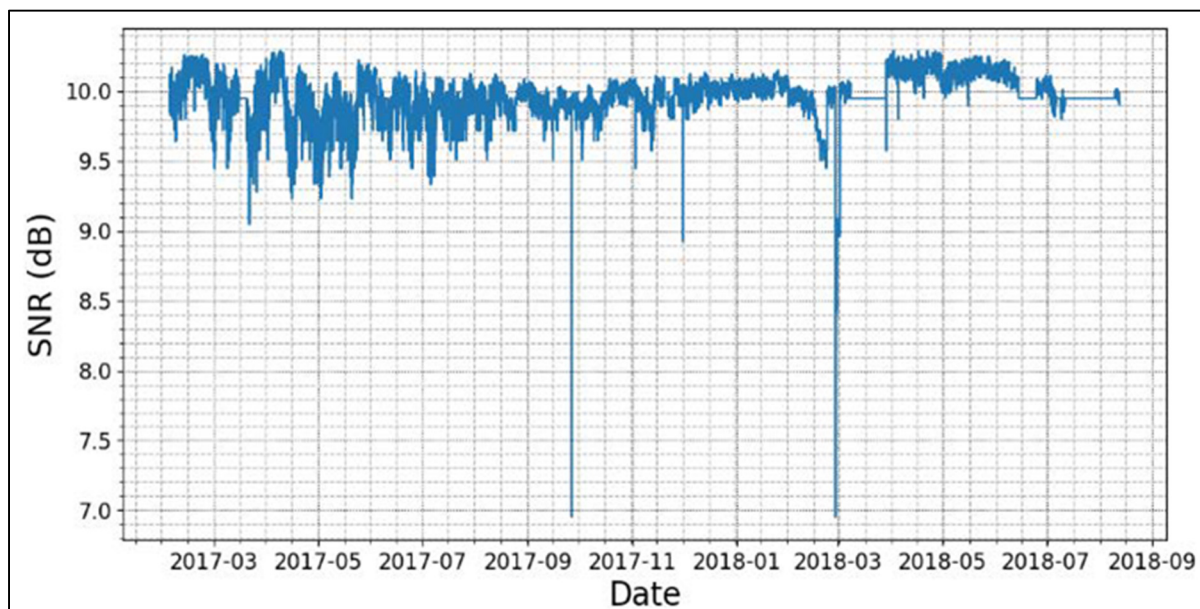


Figure 3.3 Évolution temporelle du SNR de la connexion optique B (base de données B) durant la période d'observation de 18 mois
Tirée de Aladin et al. (2020b)

La Figure 3.3 présente l'évolution temporelle du SNR de la connexion optique B (base de données B) durant la période d'observation de 18 mois. Les valeurs de SNR varient entre 6,95 dB à 10,29 dB avec une valeur moyenne de 9,96 dB et un écart-type de 0,16 dB.

3.3.2 Prétraitement des données

Les Figure 3.2 et Figure 3.3 montrent la présence de données manquantes dans les séries temporelles du SNR, à savoir 1 020 instances manquantes pour la base de données A et 7 150 instances manquantes pour la base de données B. Comme la présence de données manquantes a un impact sur les performances des modèles prédictifs, trois méthodes de traitement peuvent être utilisées pour gérer les données manquantes (Singh, Thakur et Sharma, 2016) :

- l'imputation qui consiste à remplacer les données manquantes par une valeur spécifique (moyenne, médiane, valeur maximale...),
- la suppression qui consiste à éliminer totalement ou partiellement les données manquantes,
- l'analyse qui consiste à utiliser d'autres méthodes sur lesquelles les données manquantes n'ont aucun impact.

Dans le cadre des présents travaux, c'est la méthode par imputation qui a été utilisée. Ainsi, les données manquantes de chaque série temporelle sont remplacées par la moyenne mobile des données de SNR en utilisant une fenêtre coulissante. Des tests ont été faits pour choisir la taille de la fenêtre coulissante. Les tests ont révélé que la perte d'information était plus importante lorsque la taille de la fenêtre coulissante était inférieure à 5 heures et excédait 15 heures). La taille de la fenêtre coulissante a donc été fixée à 7,5 heures.

Une fois les trous comblés dans les séries temporelles, la deuxième analyse effectuée permet de déterminer s'il existe des tendances ou des effets périodiques ou saisonniers dans les séries temporelles. Pour ce faire, deux tests ont été effectués. Le premier test est un test de stationnarité. En effet, la stationnarité permet de vérifier si une série temporelle a ses propriétés statistiques (espérance, variance, moyenne ...) dépendantes du temps ou non. Dans une série

stationnaire, aucune variation de ses propriétés statistiques n'est observée au fil du temps. Il existe plusieurs méthodes pour étudier la stationnarité d'une série temporelle. L'une de ses méthodes consiste à observer la représentation temporelle des données afin de détecter si une éventuelle tendance se dégage de celle-ci. Cette méthode cependant est très peu fiable et nécessite plusieurs vérifications afin de valider l'affirmation faite par l'observation.

Une seconde méthode est basée sur une analyse de la moyenne et de la variance de données extraites de la série temporelle. Elle consiste dans un premier temps à segmenter la série temporelle en deux ou plusieurs segments. Par la suite, la moyenne et la variance sont calculées dans chaque segment et comparées entre elles. S'il existe une différence significative entre les valeurs calculées d'un segment à l'autre, alors la série est considérée comme non stationnaire. Dans le cas contraire, la série est stationnaire. Cette méthode bien que rapide n'est pas exhaustive et peut facilement conduire à des erreurs dans l'interprétation de la comparaison. De plus, elle requiert que les données suivent une distribution gaussienne.

La dernière méthode d'étude de la stationnarité est basée sur des tests statistiques, à savoir les tests de racine unitaire (test augmenté de Dickey Fuller (*Augmented Dickey Fuller*, *ADF*), le test de Phillips-Perron, *PP*, ...) et les tests de stationnarité (test de Kwiatkowski-Phillips-Schmidt-Shin, *KPSS*, test de Leybourne et McCabe). Ces tests permettent de valider ou non l'hypothèse de stationnarité faite sur la série temporelle en comparant la valeur du test avec la valeur par défaut déterminée à partir d'un seuil de confiance. Dans le cadre de nos travaux, seuls le test ADF et le test de KPSS sont utilisés. Ainsi, pour le test ADF, l'hypothèse nulle posée H_0 suppose que la série temporelle est générée par une racine unitaire, et donc, elle n'est pas stationnaire. Comme décrit dans l'équation (3.2), si la valeur du test est supérieure à la valeur du seuil définie par le niveau de confiance, l'hypothèse nulle H_0 ne peut être rejetée et donc la série temporelle est considérée comme non stationnaire. Dans le cas contraire, si la valeur du test est inférieure à la valeur du seuil de confiance, l'hypothèse nulle H_0 est rejetée et la série est considérée comme stationnaire.

Posons H_0 : la série temporelle a une racine unitaire (3.2)

Si $test_{value} > p_{value} \rightarrow H_0$ ne peut être rejetée

Si $test_{value} < p_{value} \rightarrow H_0$ est rejetée

où :

- H_0 représente l'hypothèse nulle,
- $test_{value}$ représente la valeur statistique du test,
- p_{value} représente la valeur au seuil de confiance.

Concernant le test de KPSS, l'hypothèse nulle H_0 , quant à elle, suppose que la série temporelle est stationnaire. Ainsi, comme décrit dans l'équation (3.3), si la valeur du test est supérieure à la valeur du seuil de confiance, l'hypothèse nulle H_0 est rejetée et donc la série temporelle est non stationnaire. Dans le cas contraire, si la valeur du test est inférieure à la valeur du seuil de confiance, l'hypothèse nulle H_0 est acceptée et la série est considérée comme stationnaire.

Posons H_0 : la série temporelle a une racine unitaire (3.3)

Si $test_{value} > p_{value} \rightarrow H_0$ est rejetée

Si $test_{value} < p_{value} \rightarrow H_0$ est acceptée

où :

- H_0 représente l'hypothèse nulle,
- $test_{value}$ représente la valeur statistique du test,
- p_{value} représente la valeur au seuil de confiance.

Le Tableau 3.2 présente les résultats des analyses de stationnarité effectuées sur les deux bases de données. Les différents tests ont été effectués en utilisant la librairie *Statsmodels* de *Python*.

Tableau 3.2 Résultats des tests de stationnarité

	Base de données A		Base de données B	
	testvalue	pvalue	testvalue	pvalue
Test ADF	-4,88	-2,86	-9,42	-2,57
Test de KPSS	3,61	0,176	10,48	2
Conclusion	Non stationnaire		Non stationnaire	

Comme on peut voir dans le Tableau 3.2, les tests de stationnarité effectués révèlent que les deux séries temporelles de SNR ne sont pas stationnaires. La non-stationnarité d'une série temporelle peut avoir un impact sur les performances des algorithmes de prédiction utilisées. Ainsi, afin de rendre stationnaires les séries temporelles, il existe plusieurs méthodes de transformation appliquées aux séries, à savoir (Brockwell et Davis, 2002) :

- la différenciation consistant à soustraire de façon consécutive la valeur précédente à la valeur actuelle. En d'autres termes, la série temporelle X à l'instant t est déterminée par la formule $X_t = X_t - X_{t-1}$;
- la transformation logarithmique consistant à appliquer la fonction logarithmique à la série temporelle originale. En d'autres termes, la série temporelle X à l'instant t est déterminée par la formule $X_t = \log(X_t)$;
- la transformation puissance consistant à appliquer une puissance (carré, cubique, racine carrée ...) à la série temporelle. En d'autres termes, la série temporelle X à l'instant t est déterminée par la formule $X_t = X_t^n$, où n est le degré de la puissance.

Dans le cadre des présents travaux, la différenciation est appliquée aux deux séries temporelles. Cette méthode permet de stabiliser la moyenne de la série temporelle en supprimant les dépendances temporelles observées entre les éléments consécutifs de la série temporelle.

La deuxième analyse de tendance effectuée sur les séries temporelles est l'analyse de saisonnalité. L'analyse de saisonnalité a principalement porté sur l'analyse de la saisonnalité journalière des données. Deux méthodes ont été étudiées dans cette analyse. Notons que les tests ont été effectués en utilisant la librairie *Statsmodels* de *Python*.

La première méthode consiste à analyser les composantes des séries temporelles. Pour ce faire, les séries temporelles sont décomposées en trois composantes : la composante tendancielle, la composante saisonnière, la composante résiduelle. Cette décomposition est effectuée à l'aide de la méthode STL (*Seasonal-Trend decomposition using LoESS*) (Yaméogo et al., 2020). Par la suite, l'amplitude crête à crête des composantes saisonnières de 24 heures extraites de chaque série temporelle a été calculée. Ce calcul révèle que les séries temporelles ne comportent pas d'effets saisonniers, avec des amplitudes de saisonnalité (0,0275 dB et 0,06 dB, respectivement pour les bases de données A et B) non significatives.

La deuxième méthode a été effectuée en évaluant la force de saisonnalité des séries temporelles (Hyndman et Athanasopoulos, 2018). Celle-ci est déterminée à partir de la métrique F_s et donne une mesure de la force de la saisonnalité entre 0 (pas de saisonnalité) et 1 (forte saisonnalité). Ainsi, les valeurs de F_s (0,003 et 0,008, respectivement pour les bases de données A et B) révèlent que les séries temporelles ne comportent pas d'effets saisonniers.

Une fois l'étape de prétraitement des données finalisée, les bases de données sont divisées chacune en trois parties selon un ratio précis. Chacune de ces parties est utilisée pour la base d'apprentissage, la base de validation et la base de test.

3.4 Modèles de prédiction de performance d'une connexion optique

3.4.1 Description des modèles de prédiction de SNR

Dans cette section, des modèles univariés sont proposés pour la prédiction du SNR d'une connexion optique pour des horizons variant de 1 h à 96 h, en se basant sur un historique

d'observations du SNR durant plusieurs mois (Aladin et al., 2020a; Aladin et al., 2020b; Tremblay, Allogba et Aladin, 2019).

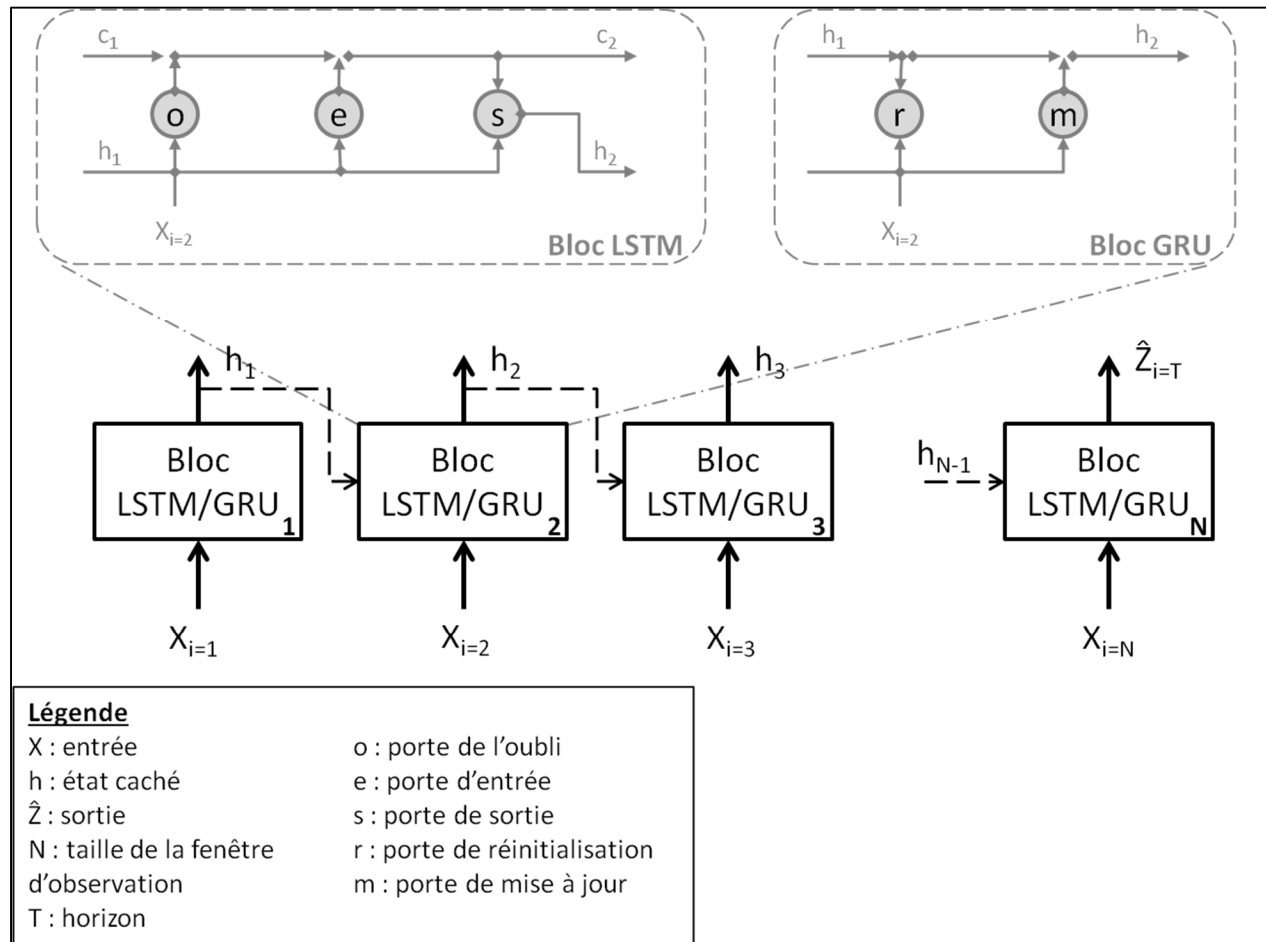


Figure 3.4 Topologie des réseaux LSTM et GRU

Deux variantes des réseaux de neurones, dont les topologies sont présentées à la Figure 3.4, sont utilisées pour l'implémentation des différents modèles :

- le LSTM univarié : comme décrit dans la Figure 3.4, le modèle LSTM utilise principalement trois structures appelées portes (*gate*), à savoir la porte de l'oubli, la porte d'entrée et la porte de sortie. La porte de l'oubli (*forget gate*) fait un filtre des informations contenues dans la mémoire des cellules (l'état caché du réseau h ainsi que l'état précédent

- du bloc c) aux instants précédents $t-1$. La porte d'entrée (*input gate*) filtre les informations reçues X à l'entrée de chaque bloc à l'instant courant t . Ces deux portes permettent ainsi d'extraire les informations pertinentes de chaque cellule afin de déterminer la sortie désirée $\hat{Z}$ à l'horizon T ;
- le GRU univarié : contrairement au LSTM univarié, le GRU univarié ne dispose que de deux portes. La porte de réinitialisation a la même fonction que la porte de l'oubli. Elle permet de contrôler le nombre d'informations de la mémoire aux instants précédents (X et h) à ignorer. La porte de mise à jour est quant à elle une combinaison de la porte d'oubli et de la porte d'entrée. Les composantes les plus importantes des données passées (X et h) sont traitées afin de déterminer la sortie à prédire ;

En d'autres termes, les données d'entrée X sont réparties en plusieurs séquences de taille N . Ces données X correspondent aux données du SNR des différentes bases de données. La taille des séquences correspond à la période d'observation utilisée pour prédire la donnée $\hat{Z}$. Pour chaque valeur de la séquence d'entrée, le modèle LSTM ou GRU met à jour leurs états internes h jusqu'aux derniers éléments de la séquence X_N . Le modèle par la suite produit une sortie $\hat{Z}$ correspondant à la prévision du SNR pour l'horizon T .

Par ailleurs, plusieurs hyper-paramètres sont à prendre en compte dans les modèles univariés LSTM et GRU, à savoir :

- le pas d'apprentissage : il contrôle la vitesse d'apprentissage du modèle. Les valeurs du pas d'apprentissage sont comprises entre 0 et 1. Ainsi, un pas d'apprentissage trop élevé peut conduire à une convergence rapide du modèle vers une solution sous-optimale alors qu'un taux de convergence trop faible conduit à une lenteur d'exécution du modèle ;
- le taux d'abandon (*dropout rate*) : il représente la probabilité d'utiliser ou d'ignorer aléatoirement un neurone du réseau pendant la phase d'entraînement. Le taux d'abandon permet donc d'éviter le sur apprentissage et a un impact sur le temps de calcul. Il est compris entre 0 et 1 ;
- le nombre de couches cachées : il représente la taille de l'état caché ;

- le nombre d'époques : correspondant au nombre d'itérations à effectuer pour l'apprentissage du modèle. Plus le pas d'apprentissage est faible, plus le nombre d'époques est important. De même, plus le pas d'apprentissage est élevé, plus le nombre d'époques est faible ;
- la fonction d'activation : elle correspond à la méthode utilisée à la fin de chaque couche du réseau pour activer les réponses de l'état caché.

3.4.2 Implémentation des modèles de prédiction du SNR

Cette section décrit comment les modèles de prédiction du SNR ont été implémentés et testés. Deux modèles de prédiction du SNR des connexions optiques A et B ont été développés : un modèle univarié LSTM construit avec la base de données A (modèle LSTM-U-1A) et un modèle univarié GRU (modèle GRU-U-1B) construit avec la base de données B. Une étude comparative des modèles univariés LSTM et GRU (LSTM-U-2A et GRU-U-2A) pour la prédiction du SNR de la base de données A est ensuite effectuée.

La première étape pour l'implémentation des modèles de prédiction a été la répartition de la base de données résultant de la phase de prétraitement en trois sous-ensembles pour la construction, la validation des hyper-paramètres et l'évaluation des modèles respectivement: l'ensemble d'apprentissage, l'ensemble de validation et l'ensemble de test.

Dans ce contexte, le choix du ratio des ensembles d'apprentissage et de test a été fait par optimisation. Notons qu'un ensemble de validation n'a pas été utilisé pour la construction du modèle LSTM-U-1A. Le ratio optimum a été établi en fixant les hyper-paramètres du modèle LSTM (nombre de couches cachées = 1, époque = 150, taux d'abandon = 0.2, taux d'apprentissage = 0,005, fonction d'activation = Adam) et en calculant le RMSE pour chacun des ratios suivants : 70/30, 80/20 et 90/10. Ces ratios correspondent respectivement à 70%, 80% et 90% de la base de données A pour l'ensemble d'apprentissage et 30%, 20% et 10% de la base de données A pour l'ensemble de test. Le ratio optimal pour le modèle LSTM-1A est

celui correspondant à la plus petite valeur de RMSE, soit 70/30. Cette répartition équivaut à 9 mois d'observation pour l'ensemble d'apprentissage et 4 mois d'observation pour l'ensemble de test sur les 13 mois d'observation de la base de données A. Pour les modèles GRU-U-1B, LSTM-U-2A et GRU-U-2A, la taille des ensembles d'apprentissage et de test a été choisie arbitrairement en se basant sur des valeurs typiques. Le ratio utilisé est de 80/20 correspondants à 80% de la base de données pour l'ensemble d'apprentissage et 20% de la base de données pour l'ensemble de test. De plus, l'ensemble de validation a également été défini en prenant 20% de l'ensemble d'apprentissage.

La seconde étape d'implémentation concerne l'optimisation des hyper-paramètres des modèles. Cette optimisation est réalisée à partir des données de l'ensemble d'apprentissage pour le modèle LSTM-U-1A et des données de l'ensemble de validation pour les modèles GRU-U-1B, LSTM-U-2A et GRU-U-2A. Il existe principalement deux techniques d'optimisation des hyper-paramètres : l'optimisation manuelle et l'optimisation automatique (Wu et al., 2019). L'optimisation manuelle consiste à tester le modèle en utilisant une valeur d'hyper-paramètre choisie à la fois. Chaque valeur est choisie aléatoirement et le modèle est évalué pour chaque valeur différente en observant une métrique donnée. Le choix définitif des hyper-paramètres est basé sur le modèle donnant les meilleures performances, sur la base de la métrique observée. Cette méthode est utile lorsque le nombre d'hyper-paramètres à optimiser n'est pas élevé. Cependant, lorsque le nombre d'hyper-paramètres à optimiser est élevé, cette méthode requiert des quantités et des temps de calcul de plus en plus importants. L'optimisation automatique, quant à elle, répond aux inconvénients de l'optimisation manuelle. Selon cette méthode, le modèle est évalué automatiquement soit en utilisant une combinaison de valeurs aléatoire (recherche aléatoire) ou non (recherche par grille) des hyper-paramètres, soit en utilisant les informations précédentes pour ajuster le choix des hyper-paramètres (recherche bayésienne). L'optimisation automatique requiert néanmoins un temps de calcul élevé. Dans le cadre des présents travaux, la méthode d'optimisation choisie est l'optimisation manuelle et la métrique utilisée pour l'évaluation des modèles et le choix final

des hyper-paramètres est le RMSE. Le Tableau 3.3 récapitule les hyper-paramètres trouvés dans la phase d'optimisation.

Tableau 3.3 Hyper-paramètres des modèles de prédiction de SNR

	LSTM-U-1A	GRU-U-1B	LSTM-U-2A	GRU-U-2A
Taux d'apprentissage	0,005	<u>0,000025</u>	<u>0,00001</u>	<u>0,00001</u>
Taux d'abandon	0,2	<u>0</u>	<u>0,2</u>	<u>0</u>
Nombre de couches cachées	<u>250</u>	<u>256</u>	<u>256</u>	<u>256</u>
Taille de la fenêtre d'observation	15 min	<u>48 h</u>	<u>48 h</u>	<u>48 h</u>
Nombre d'époques	<u>50</u>	50	5	5
Fonction d'activation	Adam – $Tanh^1$ – sigmoïde	Adam – $Tanh^1$ – sigmoïde	Adam – $Tanh^1$ – sigmoïde	Adam – $Tanh^1$ – sigmoïde

¹Tangente hyperbolique

Comme présentées dans le Tableau 3.3, pour toutes les études de cas confondues, les fonctions d'activations utilisées sont :

- la fonction d'estimation du moment adaptatif (*Adaptive moment*, *Adam*) utilisé comme solveur,
- la fonction tangente hyperbolique (*Tanh*) utilisée pour déterminer l'état des cellules,
- la fonction sigmoïdale utilisée pour calculer les portes dans chaque bloc LSTM.

En outre, pour le modèle LSTM-U-1A, seuls le nombre d'unités et le nombre d'époques ont été optimisés. Ainsi, le modèle a été construit pour des valeurs de nombre d'unités compris dans les plages (1 ; 10 ; 100 ; 200 et 250) et de nombre d'époques compris dans les plages (10 ; 50 ; 100 ; 150 et 200). Le plus petit RMSE a été trouvé pour les valeurs avec 250 pour le nombre de couches cachées et 50 pour le nombre d'époques.

Par ailleurs, afin d'évaluer les performances des modèles GRU-U-1B, LSTM-U-2A et GRU-U-2A, les hyper-paramètres à optimiser sont le pas d'apprentissage, le taux d'abandon et le nombre d'unités. En effet, ces différents hyper-paramètres ont un impact sur la rapidité d'exécution du modèle ainsi que la performance du résultat trouvé. Les modèles ont donc été évalués pour différentes valeurs d'hyper-paramètres, à savoir (0,00001 ; 0,000025 et 0,00005) pour le pas d'apprentissage, (0 et 0,2) pour le taux d'abandon et (256 et 512) pour le nombre de couches cachées. Ainsi, le nombre de couches cachées obtenu pour tous les modèles confondus est 256. Quant aux autres hyper-paramètres, pour le modèle GRU-U-1B, le plus petit RMSE a été obtenu pour les valeurs 0,000025 et 0 correspondant respectivement au taux d'apprentissage et au taux d'abandon. Pour les modèles LSTM-U-2A et GRU-U-2A, les valeurs de pas d'apprentissage et de taux d'abandon trouvées sont respectivement (0,00001 ; 0,2) et (0,00001 ; 0). Le nombre d'époques est fixé arbitrairement à 50 pour le modèle GRU-U-1B et 5 pour les modèles LSTM-U-2A et GRU-U-2A. En effet, au-delà de ces valeurs, le RMSE observé se stabilise, quelle que soit la valeur de l'hyper-paramètre.

Par ailleurs, comme présentée dans la Figure 3.4, la taille de la fenêtre d'observation intervient également dans l'implémentation du modèle. Celle-ci représente le nombre d'observations (les éléments d'entrées X) précédentes à utiliser pour prédire la valeur. Pour le modèle LSTM-U-1A, la taille de la fenêtre a été fixée à 15 minutes. Quant aux modèles GRU-U-1B, LSTM-U-2A et GRU-U-2A, la même procédure que celle de l'optimisation des hyper-paramètres a été utilisée. Le RMSE a donc été calculé pour les fenêtres d'observation de 24 h, 48 h et 96 h. La fenêtre d'observation obtenue pour tous les modèles est de 48 heures. Autrement dit, afin de déterminer la valeur à prédire à l'instant $(t + T)$, le modèle LSTM-U-1A utilise directement la

valeur précédente correspondant à l'instant $(t - 15 \text{ min})$, tandis que les modèles GRU-U-1B, LSTM-U-2A et GRU-U-2A utilisent la valeur correspondant à l'instant $(t - 48 \text{ h})$.

3.5 Évaluation de performance

Pour des fins de comparaison dans la phase d'évaluation, un modèle naïf a également été implémenté en utilisant l'ensemble de test pour chaque base de données. Ce modèle naïf consiste à attribuer à la valeur du SNR à prédire, la dernière valeur de la fenêtre d'observation. En d'autres termes, pour un horizon T , la valeur du SNR prédite à l'instant $(t + T)$ en utilisant le modèle naïf correspond à la valeur observée à l'instant $(t - 1)$. Notons que d'autres modèles de base, tels que la régression linéaire, pourraient également être utilisés comme base de comparaison.

En plus du modèle naïf, pour les modèles GRU-U-1B, LSTM-U-2A et GRU-U-2A, une variante du réseau LSTM, l'encodeur-décodeur LSTM a été implémenté dans le but de comparer ses performances avec les modèles GRU et LSTM. Celui-ci est principalement utilisé pour les tâches séquentielles. Il se compose d'un codeur LSTM qui comprime la séquence de données du SNR en un vecteur de taille fixe. La longueur de la séquence de données du SNR dépend de la taille d'observation. Aussi, à chaque pas de temps dans la séquence d'entrée, l'encodeur-décodeur LSTM met à jour sa cellule de mémoire interne jusqu'à atteindre la valeur finale désirée (valeur à prédire du SNR). Les valeurs d'hyper-paramètres ont également été optimisées en utilisant la base de données de validation. Les hyper-paramètres trouvés sont 0,0005 pour le pas d'apprentissage, 0 pour le taux d'abandon et 512 pour le nombre de couches cachées.

Les modèles, incluant le modèle naïf et le modèle encodeur-décodeur LSTM, sont donc évalués pour des horizons de 1 à 96 heures. Pour ce faire, deux métriques ont été utilisées : le RMSE et le R^2 . Ces métriques sont déterminées en effectuant une seule simulation des modèles sur l'ensemble de test. Le RMSE permet d'évaluer la précision du modèle en indiquant à quel

point les valeurs prédites sont proches des valeurs observées. Les valeurs de RMSE varient entre 0 et 1. Ainsi, plus celles-ci sont faibles, meilleure est la prédiction faite par le modèle. Le R^2 détermine la robustesse du modèle en montrant dans quelle mesure les hyper-paramètres choisis expliquent la variabilité des données à prédire. Les valeurs de R^2 sont inférieures ou égales à 1. Contrairement au RMSE, plus la valeur de R^2 est élevée, plus robuste est le modèle pour prédire les données futures. En outre, une valeur nulle du R^2 indique que le modèle se rapproche d'un modèle naïf.

3.5.1 Modèle LSTM-1A : résultats

Les performances du modèle LSTM-U-1A et du modèle naïf sont évaluées à partir du RMSE pour des horizons de prédiction variant de 2 à 24 h. Les simulations ont été réalisées dans l'environnement MATLAB. Les résultats sont montrés dans la Figure 3.5.

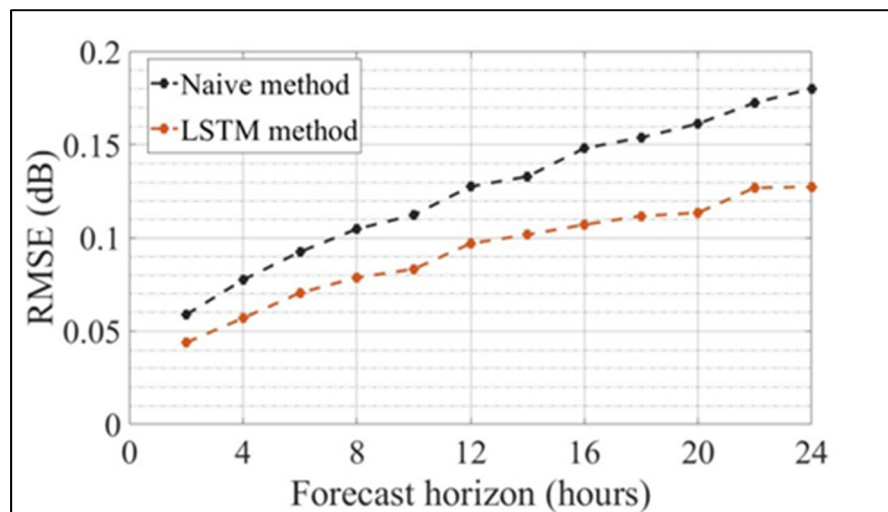


Figure 3.5 Présentation du RMSE en fonction de l'horizon pour la méthode LSTM-U-1A
Tirée de Tremblay, Allogba et Aladin (2019)

Comme on peut voir à la Figure 3.5, le RMSE augmente en fonction de l'horizon de prédiction. Cela signifie que plus l'horizon de prédiction augmente, moins bonne est la prédiction.

Néanmoins, le modèle LSTM surpasse le modèle naïf quel que soit l'horizon de prédiction, avec une amélioration maximale de 0,05 dB à 24 heures, comparativement à la méthode naïve.

La Figure 3.6 présente le SNR réel (observé) et le SNR prédit incluant l'erreur de prédiction sur deux périodes de 24 heures prises dans l'ensemble de test.

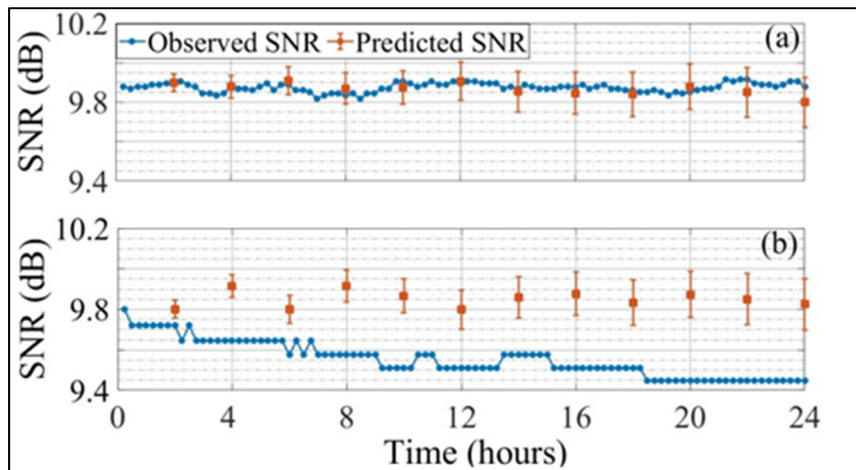


Figure 3.6 SNR prédit vs SNR observé (connexion A, LSTM-U-1A)
pour les périodes de : (a) 5 janvier ; (b) 17 janvier
Tirée de Tremblay, Allogba et Aladin (2019)

Les 12 valeurs du SNR prédites, présentées dans la Figure 3.6, ont été obtenues en utilisant 12 modèles de LSTM différents. Chacun de ces modèles a été construit en faisant varier l'horizon temporel de 2 à 24 h (avec un pas de 2 h), comme dans l'exemple présenté par Choudhury et al. (2018). Les barres d'erreur correspondent au RMSE obtenu. Le cas (a) de la Figure 3.6 présente de légères variations dans l'observation du SNR, tandis que le cas (b) présente une variation soudaine de 0.2 dB dans l'observation du SNR. Figure 3.6 montre que notre modèle LSTM construit et adapté pour chaque horizon arrive à prédire les données du SNR avec de faibles erreurs. Cependant, dans le cas de la baisse soudaine du SNR, la prédiction du modèle devient plus faible et le taux d'erreur plus élevé. Le modèle ne parvient pas à prédire la baisse du SNR.

Le temps d'exécution du modèle LSTM-U-1A est d'environ 3,3 ms sur un système Windows 10 avec un processeur Intel® Core™ i5-3210M à 2,5 GHz et une mémoire RAM de 8 Go.

3.5.2 Modèle GRU-U-1B : résultats

Le modèle GRU-U-1B a été implémenté en utilisant le package *Keras* dans Python 3. Le RMSE a également été calculé pour des horizons de prédiction variant de 1 à 48 heures, afin d'évaluer et de comparer les performances des modèles GRU, encodeur-décodeur LSTM et du modèle naïf. Ainsi, la Figure 3.7 montre l'évolution du RMSE en fonction de l'horizon.

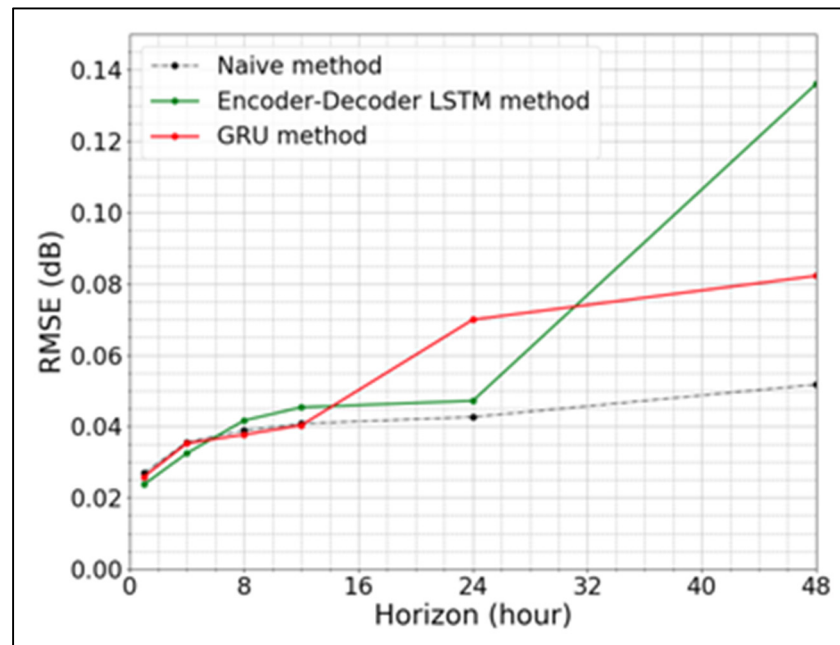


Figure 3.7 RMSE en fonction de l'horizon de prédiction pour les modèles GRU-U-1B, encodeur-décodeur LSTM et naïf
Tiré de Aladin et al. (2020b)

Pour des horizons variant de 1 h à 12 h, la performance du modèle GRU est similaire à celle du modèle naïf, avec un avantage ne dépassant pas 0,001 dB en RMSE. La performance du modèle GRU se dégrade rapidement pour des horizons variant entre 12 et 24 heures. La bonne

performance du modèle naïf résulte des variations relativement faibles du SNR observées dans l'ensemble de données de test.

Pour des horizons variant de 6 h à 12 h et plus grand que 24 h, le modèle GRU surpasse le modèle encodeur-décodeur LSTM avec un avantage allant jusqu'à 0,048 dB. De plus, plus l'horizon est lointain, plus le modèle GRU se stabilise, contrairement au modèle encodeur-décodeur LSTM.

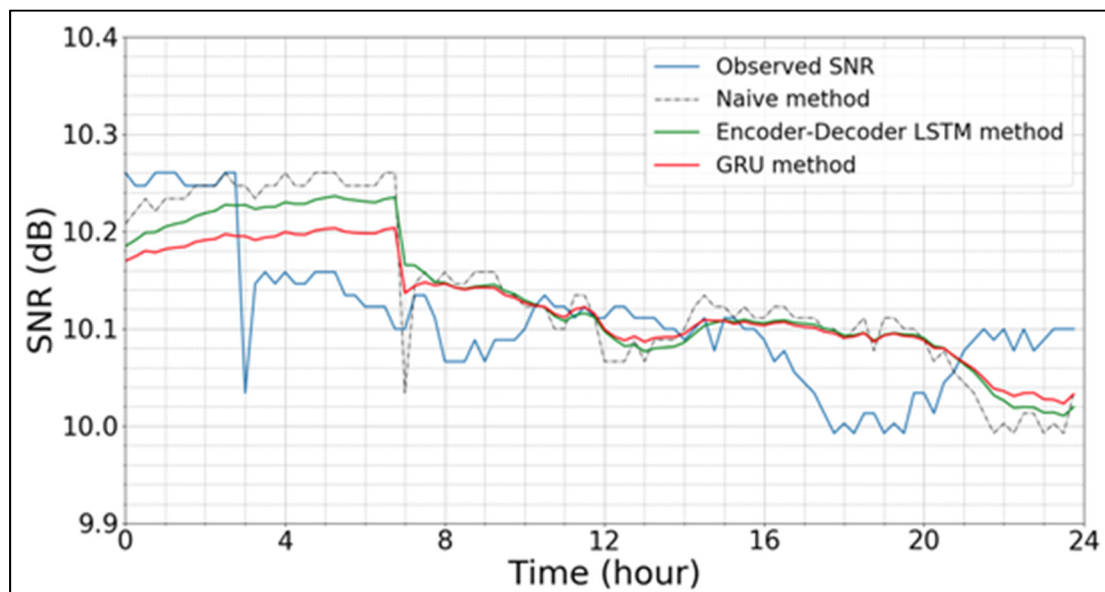


Figure 3.8 SNR prédit vs SNR observé (connexion B, GRU-U-1B) pour la période d'observation du 29 avril 2018 au 30 avril 2018
Tirée de Aladín et al. (2020b)

La Figure 3.8 montre le SNR prédit par rapport au SNR observé sur une période de 24 heures du 29 avril 2018 à 1 h 15 au 30 avril 2018 à 1 h 15, pour un horizon de prévision de 4 heures. Celle-ci confirme les résultats du RMSE pour des courts horizons, à savoir que le modèle GRU construit performe légèrement mieux que la méthode naïve et le modèle encodeur-décodeur LSTM.

Le temps d'exécution du modèle GRU-U-1B est d'environ 0,228 s quel que soit l'horizon de prédiction. Le modèle a été construit sur un système Windows 10 avec un processeur Intel® Core™ i5-8600K à 3.6 GHz.

3.5.3 Modèles LSTM-U-2A et GRU-U-2A : résultats

Les modèles LSTM-U-2A et GRU-U-2A ont été construits en utilisant le package *Keras* dans Python 3. En plus du RMSE, le R^2 a également été calculé afin d'évaluer la performance des modèles. Les horizons observés dans la phase d'évaluation varient de 1 h à 96 h. La Figure 3.9 présente respectivement le RMSE et le R^2 en fonction de l'horizon pour tous les modèles.

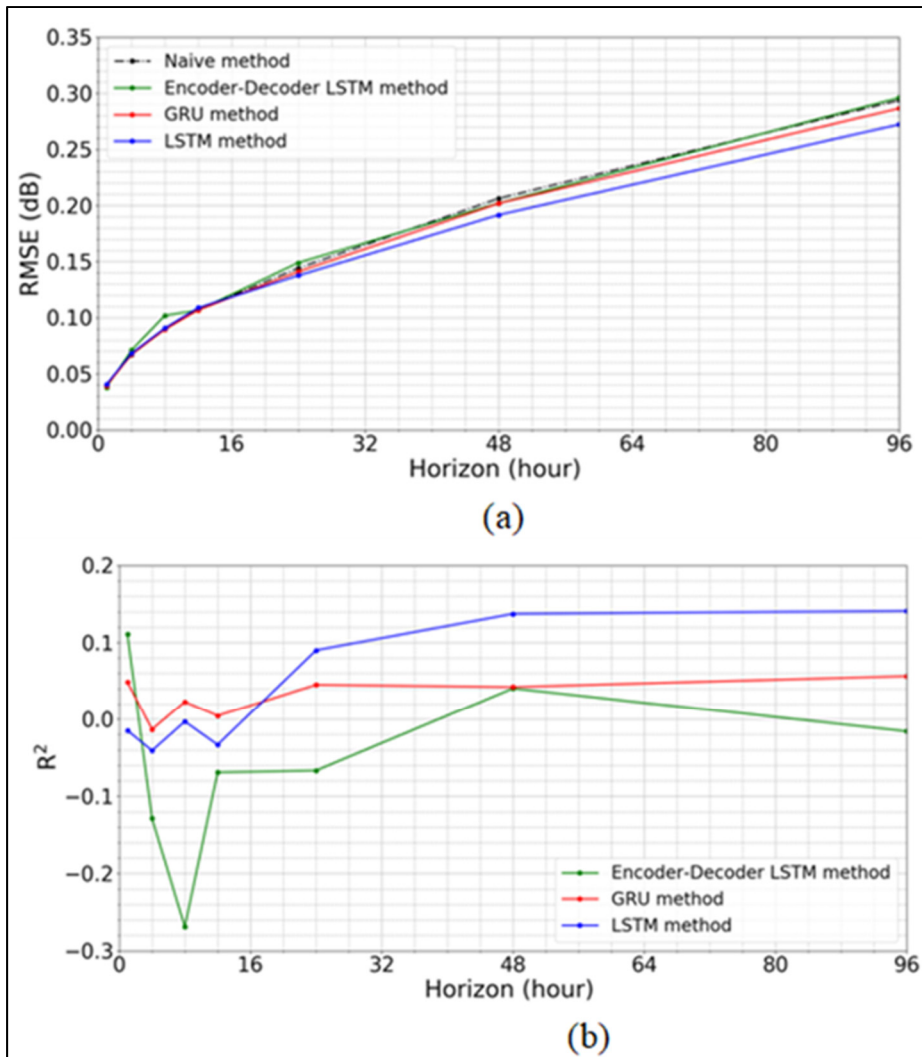


Figure 3.9 Performance des modèles LSTM-U-2A et GRU-U-2A :
 (a) RMSE en fonction de l'horizon ; (b) R^2 en fonction de l'horizon
 Tirée de Aladin et al. (2020a)

Tout comme les études de cas précédents, le RMSE augmente en fonction de l'horizon pour tous les modèles. De plus, le RMSE des modèles GRU, LSTM, encodeur-décodeur LSTM et naïf sont similaires pour les horizons allant de 1 h à 24 h. Pour les horizons de 24 h à 96 h, les modèles LSTM et GRU surpassent progressivement le modèle naïf et le modèle encodeur-décodeur LSTM avec des améliorations maximales de 0,022 dB et 0,012 dB à 96, respectivement pour le modèle LSTM et le modèle GRU.

Concernant le R^2 présenté à la Figure 3.9(b), les modèles GRU et LSTM donnent globalement de meilleures performances que la méthode naïve, principalement pour des horizons élevés (24 h, 48 h et 96 h). En effet, en regardant plus en détail, le R^2 du modèle LSTM commence par -0,0148 à 1 heure et s'améliore progressivement avec un horizon de prévision plus lointain se terminant par une valeur de 0,1415 à 96 heures. Cela signifie donc que le modèle LSTM surpasse une méthode basique pour les horizons de 24, 48 et 96 heures. Quant au modèle GRU, les valeurs R^2 sont majoritairement positives allant de -0,0133 (prédiction sur 4 heures) à 0,0563 (prédiction sur 96 heures). Cela indique que le modèle GRU obtient une prédiction légèrement plus précise que la méthode naïve, sauf pour un horizon de 4 heures. De plus, les valeurs R^2 du modèle GRU proches de zéro montrent qu'il se comporte de manière similaire à la méthode naïve en prédisant une valeur SNR proche de la valeur la plus récente dans la fenêtre d'observation. Néanmoins, la Figure 3.9(b) présente des valeurs de R^2 s'améliorant en fonction de l'horizon. Cela laisse sous-entendre que les modèles deviennent plus performants à mesure que l'horizon augmente. Cette observation est en contradiction avec les résultats présentés à la Figure 3.9(a) où on peut voir une augmentation des valeurs de RMSE en fonction de l'horizon, traduisant une réduction des performances des modèles en fonction de l'horizon. Une étude plus approfondie sur la métrique R^2 devra donc être effectuée dans le but de fournir plus d'explications sur ces résultats contradictoires.

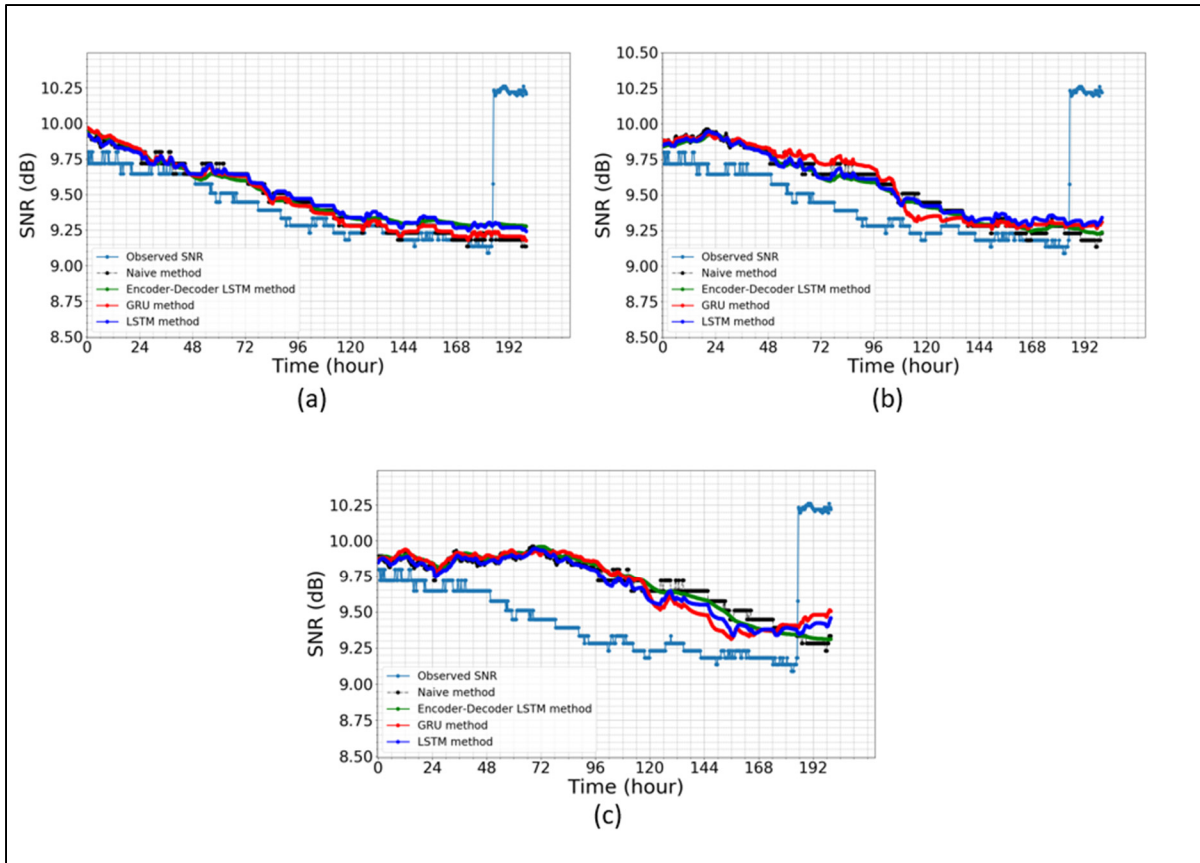


Figure 3.10 SNR prédit vs SNR observé (connexion A, LSTM-U-2A et GRU-U-2A) :
 (a) horizon de 24 heures ; (b) horizon de 48 heures ; (c) horizon de 96 heures

La Figure 3.10 montre le SNR observé et le SNR réel pendant trois périodes de 192 heures dans l'ensemble de données de test pour la période du 14 février 2018 au 22 février 2018 pour trois horizons de prévision différents (24 heures, 48 heures et 96 heures). L'objectif de cette figure est de déterminer si les modèles implémentés parviennent à prédire les données d'OSNR pour des échelles de temps de plusieurs heures. Ainsi, sur la Figure 3.10 (a) et la Figure 3.10 (b), nous voyons que le LSTM se comporte de la même manière que la méthode naïve pour l'horizon de prévision de 24 heures et 48 heures. Pour les prévisions à plus long terme, sur la Figure 3.10 (c), la courbe de prédiction LSTM est légèrement plus proche de la courbe réelle que les deux autres modèles. Cela pourrait expliquer en partie pourquoi le modèle LSTM a le RMSE le plus bas parmi les techniques étudiées. Cependant, sur la Figure 3.10(c), la courbe de prédiction du modèle GRU semble adopter une forme différente de la courbe naïve après

120 heures (5 jours), devenant légèrement plus proche des données de SNR réelles. En outre, les courbes de prédiction du modèle GRU pour les horizons 24 heures et 48 heures semblent être les plus proches de la courbe des valeurs réelles observées.

De façon générale, le modèle LSTM implémenté se présente comme le meilleur modèle comparativement au modèle GRU et naïf pour les horizons de prédiction plus élevés (24, 48 et 96 h). En effet il donne des valeurs de R^2 plus élevées ainsi que des valeurs de RMSE plus faibles à ces horizons. Pour des horizons compris entre 1 h et 12 h, le modèle LSTM ne surpasse pas le modèle naïf. Par contre, en considérant tous les horizons, le modèle GRU se positionne comme le meilleur modèle. En effet, il possède des valeurs de R^2 positives et un RMSE plus faible que le modèle naïf, quel que soit l'horizon observé.

Les modèles ont été exécutés sur un système avec un processeur Intel® Core™ i5-8600K 3,6 GHz avec 16 Go de mémoire RAM et un GPU GTX 970. Le temps de calcul des modèles LSTM-U-2A et GRU-U-2A est resté stable à environ 47,7 ms et 40,8 ms respectivement.

3.6 Conclusion

Les travaux présentés dans ce chapitre ont permis d'explorer deux algorithmes de prédiction de SNR de deux connexions optiques, à savoir le LSTM et le GRU, selon trois scénarios différents. Par ailleurs, ces algorithmes de prédiction ont été introduits pour la première fois dans le domaine des réseaux optiques afin de prédire les performances de connexions optiques déjà établies.

En effet, dans le premier scénario, le modèle LSTM-U-1A a été implémenté pour la prédiction du SNR en se basant sur des données de terrain d'une connexion optique recueillies sur une période 13 mois. Ce modèle avait pour but de montrer le potentiel des techniques d'apprentissage machine dans la prédiction des performances des liaisons optiques déjà

établies. La meilleure performance a été obtenue pour un horizon de prévision de 24 h pour le LSTM-U-1A par rapport à une méthode naïve, avec une différence de 0,05 dB.

Le deuxième scénario a consisté à implémenter le modèle GRU-U-1B en utilisant une base de données de SNR collectée sur une connexion optique déjà établie pendant une période de 18 mois. Ce modèle avait pour but, tout comme l'étude de cas précédent, de prédire les données de SNR, cette fois-ci pour des horizons allant jusqu'à 48 heures. Les performances observées ont permis de montrer que, pour la liaison étudiée, le modèle GRU-U-1B performe mieux qu'un modèle naïf pour des horizons de prédiction inférieurs à 12 heures.

Finalement, le troisième scénario a permis de comparer les modèles GRU et LSTM (appelés respectivement ici GRU-U-2A et LSTM-U-2A) en utilisant les données de la connexion optique A. De nouveaux hyper-paramètres ont été pris en compte dans l'implémentation de ces modèles. Les performances obtenues ont montré que les modèles LSTM-U-2A et GRU-U-2A donnent de meilleures performances que le modèle naïf et le modèle encodeur-décodeur LSTM pour des horizons de 1 à 96 heures.

Les modèles LSTM et GRU proposés dans ce chapitre sont de type univarié. Dans le prochain chapitre, des modèles multivariés utilisant des données supplémentaires sont proposés comme moyen d'augmenter la performance en prédiction.

CHAPITRE 4

MODÈLES MULTIVARIÉS POUR LA PRÉDICTION DE PERFORMANCE D'UNE CONNEXION OPTIQUE

4.1 Introduction

Un moyen potentiel de réduire l'impact des dégradations et des pannes de performances au niveau des liaisons consiste à mettre en œuvre des outils de prédiction ou des classificateurs de performance basés sur les algorithmes d'apprentissage machine. Dans ce contexte, deux nouveaux modèles LSTM et GRU multivariés sont proposés pour la prédiction de la qualité de transmission d'une connexion optique, dans la continuité des travaux présentés dans Aladin et al. (2020a). Les modèles sont entraînés avec une base de données, appelée ici KB-1, composée des données historiques de BER présentées dans la section 3.3, de la puissance optique au récepteur (P_{RX}) et du délai de groupe différentiel (DGD) d'une des connexions du réseau CANARIE (appelée ici *lightpath-1*), ainsi que de la température extérieure et de la période du jour. De plus, une analyse des différents attributs précités est également réalisée afin d'étudier leur impact sur la performance des modèles multivariés implémentés.

Par ailleurs, le potentiel de l'apprentissage par transfert est exploré en testant les modèles multivariés précédemment entraînés avec les données de la connexion *lightpath-1* sur deux ensembles de données différents, KB-2 et KB-3, issus de deux autres connexions optiques du réseau CANARIE (respectivement *lightpath-2* et *lightpath-3*). Notons que la connexion optique *lightpath-2* est transportée dans la même fibre optique que la connexion optique *lightpath-1* mais dans une portion de 600 km de la route de 1300 km dans laquelle la connexion optique *lightpath-1* est transportée. Contrairement à *lightpath-2*, la connexion optique *lightpath-3* est déployée sur la même route mais dans la direction opposée à la connexion optique *lightpath-1*, donc dans une fibre optique différente. L'objectif à travers cette nouvelle approche est d'évaluer la possibilité de prédire le SNR de connexions optiques provenant de domaines différents ayant à la fois des configurations différentes (direction opposées, fibre optique

différente) et des configurations identiques (même direction, portion de la même route) du domaine source.

Les travaux présentés dans ce chapitre font l'objet d'un manuscrit : S. Allogba, S. Aladin and C. Tremblay. 2021. « Multivariate Machine Learning Models for Short-Term Lightpath Performance Forecast », in *Journal of Lightwave Technology*, doi : 10.1109/JLT.2021.3110513.

Le reste du chapitre est organisé comme suit. L'objectif de ce chapitre est présenté dans la section 4.2 à partir d'une description de la problématique rencontrée dans la prédiction des données de SNR. Les bases de données KB-1, KB-2 et KB-3 issues respectivement de chacune des connexions optiques à l'étude (lightpath-1, lightpath-2 et lightpath-3) et les étapes de prétraitement des données sont décrites dans la section 4.3. L'implémentation des modèles de prédiction multivariés et la présentation des résultats de performance sont effectuées dans la section 4.4. L'impact des attributs est étudié dans la section 4.5. Enfin, l'évolutivité des modèles à travers le concept d'apprentissage par transfert pour la prédiction du SNR des bases de données KB-2 et KB-3 est présentée et discutée dans la section 4.6.

4.2 Objectif

En plus d'être un meilleur support pour transporter des signaux optiques sur des distances extrêmement longues, la fibre optique est aussi un bon capteur de température et de contrainte. Les performances des connexions optiques peuvent présenter des variations saisonnières et des dégradations dues aux conditions météorologiques, aux conditions du sol (en présence de glace par exemple) et à une mauvaise manipulation des câbles. Le vieillissement des câbles à fibres optiques peut également entraîner une augmentation des pertes au fil du temps. Dans ce contexte, la prédiction précise des performances des connexions optiques, avant ou après leur établissement, peut s'avérer une tâche complexe. Les modèles actuels dépendent fortement de la connaissance des paramètres du système de transmission et de la liaison optique et ne

prennent pas en compte les facteurs externes tels que les conditions météorologiques, les caractéristiques extérieures de déploiement et le vieillissement des composants.

Les performances des connexions optiques peuvent également être très stables dans le temps. Dans un tel cas, la prédiction des performances à court terme est une tâche très simple qui peut être gérée sans avoir besoin de méthodes et d'outils complexes. Toutefois, certaines connexions optiques présentent un comportement beaucoup plus dynamique qui ne peut être expliqué ou prédit par des modèles théoriques. La Figure 3.2 montre un bon exemple de connexion optique dynamique. Sur cette figure, on peut voir l'évolution du SNR pour lightpath-1 se produisant dans le réseau de production de CANARIE sur une période de 13 mois. Le SNR présente des dégradations du SNR pouvant atteindre plusieurs dB au printemps et en été, ainsi que plusieurs baisses et augmentations du SNR en hiver. Ces variations de SNR, qui durent de quelques jours à plusieurs mois, ne peuvent s'expliquer par des opérations de commutation ou d'injection/extraction de canal. En outre, plusieurs études antérieures des données de terrain collectées dans des liaisons aériennes et enfouies de différents réseaux en Amérique du Nord ont révélé différents motifs dans les séries chronologiques de SNR des connexions optiques : des variations saisonnières (c'est-à-dire un SNR plus faible pendant la saison hivernale), des variations quotidiennes des pertes dépendantes en polarisation (PDL) et une activité PDL beaucoup plus élevée pendant les jours de semaine par rapport aux fins de semaine (Michel, 2016; Tremblay et al., 2017).

Le premier objectif de ce chapitre est donc de tirer parti de l'apprentissage machine pour capturer des motifs complexes dans les séries chronologiques de SNR et de les exploiter pour la prédiction de SNR d'une connexion optique.

Les techniques d'apprentissage machine ont été explorées dans les applications de communication optique au cours des dernières années. Cependant, en raison de la rareté des données de terrain, moins d'attention a été accordée à l'implémentation d'outils de prédiction de performance d'une connexion optique déjà établie. Les quelques modèles proposés à ce jour

sont des variantes des réseaux de neurones formés uniquement avec des données historiques de BER. C'est le cas des travaux présentés dans les chapitres précédents qui proposent les algorithmes GRU et LSTM pour prédire le rapport signal sur bruit (SNR) d'un trajet optique existant sur des horizons pouvant aller jusqu'à 96 heures. Plus récemment, un réseau neuronal convolutif (*Convolutional Neural Network, CNN*) a été proposé pour prédire la performance d'un trajet optique en utilisant des données de terrain. Celui-ci capturait les changements temporels du SNR et les prédisait correctement sur des horizons allant jusqu'à 24 heures (Mezni et al., 2020).

En résumé, ces modèles de prédiction sont basés sur les réseaux de neurones univariés, utilisant seulement des données historiques de SNR d'une seule connexion optique pour prédire le SNR de cette connexion sur un horizon de quelques jours. De plus, dans l'éventualité où aucun motif n'est observable dans la série temporelle du SNR, il deviendrait difficile de prédire les données de SNR en ne tenant compte que de ses variations temporelles. En effet, la prédiction des modèles prédictifs se rapprocherait le plus possible d'une prédiction faite avec un modèle naïf comme observé dans les travaux précédents.

D'autre part, l'approche d'apprentissage par transfert a été introduite dans le cadre de la prédiction et de l'estimation de la qualité de transmission d'une connexion optique non établie (Marino et al., 2020; Mo et al., 2018b; Yu et al., 2019). Dans les travaux de Mo et al. (2018b); Yu et al. (2019), un modèle d'apprentissage par transfert basé sur l'ANN a été implémenté et entraîné à partir d'une base de données source composée de données synthétiques. Ce modèle par la suite a été utilisé pour prédire le facteur Q de différents systèmes optiques composés de métriques de performances ayant des distributions similaires à la base de données source. Marino et al. (2020), quant à eux, appliquent l'approche d'apprentissage par transfert sur un modèle SVM pour estimer la qualité de transmission de nouvelles connexions optiques, toutes issues de topologies de réseau différentes. En outre, le modèle SVM utilisé dans leur approche a été entraîné à partir de données sources d'une connexion optique issue d'une topologie différente que les précédentes, mais avec le même type de fibre et les mêmes équipements utilisés dans le réseau.

Dans ce chapitre, les contributions sont donc organisées comme suit. Tout d'abord, des modèles RNN profonds multivariés sont utilisés pour tirer parti des fonctionnalités disponibles dans les ensembles de données pour les prédictions de SNR à court terme. Ensuite, en utilisant le test de corrélation de Pearson, les attributs pertinents sont sélectionnés afin d'améliorer les performances et la vitesse des modèles prédictifs. Enfin, l'apprentissage par transfert est expérimenté sur deux ensembles de données de taille réduite.

4.3 Collecte et prétraitement des données

4.3.1 Présentation des bases de données

Les trois bases de données à l'étude (KB-1, KB-2 et KB-3) sont construites à partir des métriques de performance de trois connexions optiques établies sur une portion du réseau CANARIE. Cette portion est présentée dans la Figure 4.1.

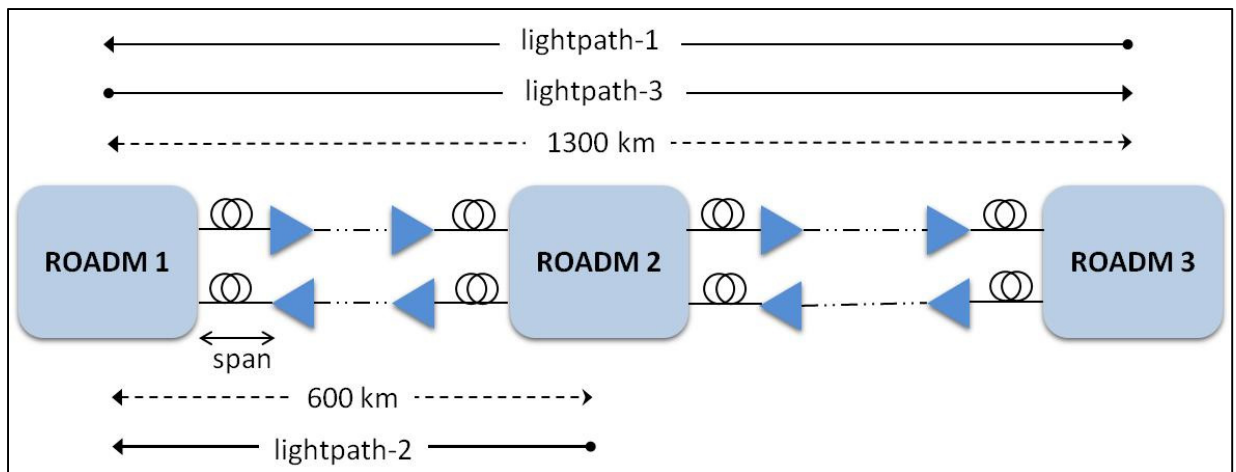


Figure 4.1 Topologie de la portion du réseau à l'étude (distances approximatives)
Adaptée de Allogba, Aladin et Tremblay (2021)

Comme décrit dans la Figure 4.1, la portion de réseau à l'étude comprend 3 nœuds ROADM interconnectés entre eux par des liaisons bidirectionnelles composées de plusieurs

amplificateurs avec un espacement moyen de 95 km. De plus, le trafic est acheminé en utilisant un nombre inconnu de canaux à 100 Gbit/s et 40 Gbit/s. Lightpath-1 et lightpath-2 sont donc issues d'une liaison bidirectionnelle d'environ 1300 km de la portion du réseau entre ROADM1 et ROADM3 avec un canal PM-QPSK à 100 Gbit/s. Lightpath-2, quant à elle, est un canal à 40 Gbit/s transporté dans une sous-section de la précédente liaison d'environ 600 km située entre ROADM1 et ROADM2.

Tableau 4.1 Caractéristiques des trois bases de données de terrain issues de trois connexions optiques utilisées pour les modèles de prédiction de performance

	KB-1	KB-2	KB-3
Connexion	Lightpath-1	Lightpath-2	Lightpath-3
Données extraites	BER DGD P_{RX}	BER P_{RX}	BER P_{RX}
Attributs	SNR DGD P_{RX} $T_{\text{extérieure}} (T_{\text{ext}})$ Période de jour	SNR P_{RX}	SNR P_{RX}
Nombre d'éléments par attributs	38 203	16 000	17 832
Données manquantes	1 020	704	711
Période d'observation	13 mois	5 mois	5 mois

Le Tableau 4.1 récapitule les informations concernant les connexions et bases de données utilisées. Premièrement, les bases de données à l'étude sont construites à partir des métriques de performance collectées pour chacune des connexions présentées dans la Figure 4.1. Ainsi,

la base de données KB-1 comprend des données de SNR, calculées à partir du pré-FEC BER, la puissance optique du canal P_{RX} et le DGD, tandis que les bases de données KB-2 et KB-3 sont composées des données de SNR et de puissance optique du canal P_{RX} .

La Figure 4.2 présente la probabilité de distribution des données de SNR des 3 connexions optiques.

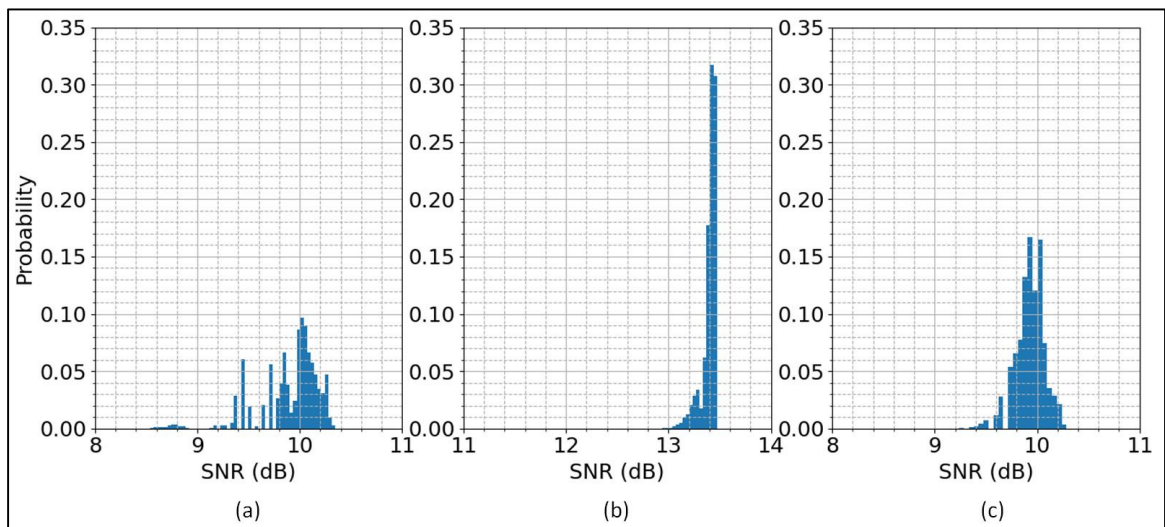


Figure 4.2 Probabilité de distribution du SNR pour : (a) lightpath-1 ; (b) lightpath-2 ; (c) lightpath-3

Tirée de Allogba, Aladin et Tremblay (2021)

En outre, chacune de ces métriques a été recueillie à la fréquence d'échantillonnage de 15 minutes durant une période de 13 mois (de février 2017 à mars 2018) pour lightpath-1, 5 mois (de février 2017 à juillet 2017) pour lightpath-2 et 18 mois (de février 2017 à août 2018) pour lightpath-3.

Rappelons que les connexions lightpath-1 et lightpath-3 sont identifiées dans la section 3.3 comme connexion A et connexion B. Dans le cadre de cette étude, nous considérons la même période d'observation pour lightpath-2 et lightpath-3. La base de données KB-1 comprend donc 38 203 éléments de chaque donnée extraite, tandis que KB-2 et KB-3 comprennent

respectivement 16 000 et 17 832 éléments de chaque donnée extraite. Il existe également des données manquantes dans chaque métrique extraite, à savoir 1 020 éléments pour KB-1, 704 éléments pour KB-2 et 711 éléments pour KB-3. Ces données manquantes pouvant affecter les performances des modèles de prédiction (Singh, Thakur et Sharma, 2016), tout comme dans la section 3.3.2, elles ont été remplacées par une moyenne mobile avec pour fenêtre fixée à 7,5 heures.

Par ailleurs, en plus des métriques de performance extraites de la connexion optique lightpath-1, deux facteurs externes sont pris en compte dans la base de données KB-1. Ces facteurs ont été choisis en fonction des résultats de nos travaux précédents (Allogba et Tremblay, 2018). Le premier facteur externe est la température horaire extérieure (T_{ext}) collectée durant la période d'observation (Canada, 2011). Celle-ci varie de -35°C à 34.4°C avec une valeur moyenne de $2,64^{\circ}\text{C}$ et un écart type de $14,53^{\circ}\text{C}$ durant la période d'observation. Le deuxième facteur externe est l'activité humaine caractérisée par quatre périodes de la journée : matinée (5 h – 12 h), après-midi (12 h – 18 h), soirée (18 h – 22 h) et nuit (22 h – 5 h).

4.3.2 Analyse statistique

Le Tableau 4.2 résume les valeurs statistiques des métriques de performance pour chacune des bases de données.

Tableau 4.2 Résumé statistique des données de SNR, P_{RX} et DGD pour chaque connexion

	Lightpath-1	Lightpath-2	Lightpath-3
--	-------------	-------------	-------------

SNR (dB)			
min	7,96	11,12	9,04
max	10,34	13,47	10,28
moyenne	9,90	13,39	9,91
sdev ¹	0,29	0,07	0,17
P_{RX} (dBm)			
min	-6,69	-8,30	-7,5
max	-4,00	-6,19	-5,5
moyenne	-5,45	-7,06	-6,15
sdev ¹	0,49	0,33	0,37
DGD (ps)			
min	2		
max	10	N/A	N/A
moyenne	5,66		
sdev ¹	1,14		

¹ Déviation standard

Ainsi pour la base de données KB-1, la série temporelle de SNR comprend 38 023 échantillons de SNR allant de 7,96 dB à 10,34 dB avec une valeur moyenne de 9,90 dB et un écart-type de 0,29 dB. De même, la série temporelle de puissance comprend également 38 023 échantillons de puissance variant de -6,69 dBm à -4 dBm avec une valeur moyenne de -5,45 dBm et un écart-type de 0,49 dBm. Enfin, la série temporelle de DGD comprend 38 023 échantillons de DGD variant de 2 ps à 10 ps avec une valeur moyenne de 5,66 ps et un écart-type de 1,14 ps. Pour la base de données KB-2, la série temporelle de SNR comprend 16 000 échantillons de SNR allant de 11,12 dB à 13,47 dB avec une valeur moyenne de 13,39 dB et un écart-type de 0,07 dB. De même, la série temporelle de puissance comprend 16 000 échantillons de puissance variant de -8,30 dBm à -6,19 dBm avec une valeur moyenne de -7,06 dBm et un écart-type de 0,33 dBm. Enfin, pour la base de données KB-3, la série temporelle de SNR comprend 17 832 échantillons de SNR allant de 9,04 dB à 10,28 dB avec une valeur moyenne

de 9,91 dB et un écart-type de 0,17 dB. De même, la série temporelle de puissance comprend 17 832 échantillons de puissance variant de -7,50 dBm à -5,50 dBm avec une valeur moyenne de -6,15 dBm et un écart-type de 0,37 dBm. Notons que la variabilité des données de SNR pour lightpath-2 et lightpath-3 est inférieure à celle des données de SNR pour lightpath-1, avec une variation maximale respective sur l'ensemble de la période d'observation de 0,08 dB et 0,39 dB pour lightpath-2 et lightpath-3, comparativement à 1,37 dB pour lightpath-1.

Les analyses statistiques effectuées sur les données sont les mêmes que celles présentées dans la section 3.3.2 et ont été exécutées en utilisant la librairie *statsmodels* dans Python 3. Celles-ci seront reprises à titre de rappel pour la base de données KB-1. Ainsi, le test de saisonnalité a révélé de faibles effets de saisonnalités avec des amplitudes crête à crête de 0,03 dB et 0,088 dB pour les composantes saisonnières de 24 heures et 7 jours. De même, concernant les données de SNR pour les bases de données KB-2 et KB-3, aucune composante saisonnière significative n'a été trouvée (amplitude crête à crête de 0,06 dB pour la composante saisonnière de 24 heures). L'ajout des attributs supplémentaires aux modèles prédictifs répond à la non-saisonnalité des données de SNR en apportant des informations supplémentaires pour la prédiction du SNR.

Les tests de stationnarité ont montré que la série temporelle de SNR de KB-1 n'est pas stationnaire (valeurs de test de -4,88 et 3,61 inférieures aux valeurs critiques de -2,86 et 0,176, respectivement pour le test ADF et KPSS). Concernant les bases de données KB-2 et KB-3, les séries temporelles de SNR ne sont également pas stationnaires. En effet, les valeurs de test pour le test ADF trouvées (-4,25 et -6,86 respectivement pour KB-2 et KB-3) sont inférieures à la valeur critique -2,86. De même, les valeurs de test pour le test KPSS (4,93 et 0,99 respectivement pour KB-2 et KB-3) sont supérieures à la valeur critique 0,146. Comme vu dans la section 3.3.2, il existe différentes méthodes pour rendre les séries chronologiques stationnaires. Cependant, avec les méthodes utilisant les algorithmes de réseaux de neurones, il n'est pas nécessaire de les rendre stationnaires (Aladin et al., 2020b; Tremblay, Allogba et Aladin, 2019; Wang et al., 2019).

Les sections suivantes présentent l'implémentation et l'optimisation des modèles multivariés.

4.4 Modèles de prédiction multivariés

Cette section est subdivisée en deux parties. La première partie présente les étapes d'implémentation des modèles multivariés LSTM-M-1 et GRU-M-1, construits à partir de la base de données KB-1. La seconde partie évalue les performances des modèles

4.4.1 Présentation des modèles multivariés et optimisation des hyper-paramètres

Les modèles multivariés considérés pour la prédiction du SNR à court et à long terme sont dérivés des algorithmes LSTM et GRU. Ils diffèrent des modèles univariés en prenant en compte, en plus des valeurs précédentes de SNR, d'autres données d'entrées pour prédire les valeurs de SNR. Ainsi, comme décrits dans la Figure 3.4 de la section 3.4.1, les modèles multivariés LSTM-M-1 et GRU-M-1 combinent la valeur observée et d'autres attributs comme entrées, contrairement aux modèles univariés qui utilisent uniquement la valeur observée comme entrée dans le modèle. En d'autres termes, afin de prédire le SNR d'une connexion optique à l'horizon T , les modèles multivariés LSTM et GRU prennent en entrée la puissance, le DGD, la température, la période du jour et les valeurs de SNR observées sur une période d'observation de N heures.

Par ailleurs, la base de données KB-1 est divisée en ensemble d'apprentissage, de validation et de test selon les ratios respectifs de 0,64 ; 0,16 et 0,20. Les modèles LSTM-M-1 et GRU-M-1 ont également été implémentés en utilisant le package *Keras* de Python 3. Les hyper-paramètres des modèles, à savoir le pas d'apprentissage, le taux d'abandon et la taille de la couche cachée ont été optimisés en utilisant l'ensemble de test avec les valeurs respectives suivantes (0,00001 ; 0,000025 ; 0,00005), (0 ; 0,2 ; 0,5 ; 0,8 ; 1) et (50 ; 150 ; 256 ; 512).

L'erreur quadratique moyenne (RMSE) a été observée comme critère pour le choix des hyper-paramètres.

Outre les hyper-paramètres, la taille de la fenêtre qui correspond à la période d'observation N dans la séquence d'entrées utilisée pour prédire le SNR a également été optimisée. Afin de déterminer la meilleure taille de la fenêtre, les valeurs (1 ; 4 ; 8 ; 12 ; 24 ; 48 heures) ont été testées avec comme critère de choix le RMSE. Le Tableau 4.3 montre l'ensemble des hyper-paramètres résultants de la phase d'optimisation.

Tableau 4.3 Présentation des hyper-paramètres des modèles

	Modèles multivariés		Modèles univariés	
	LSTM-M-1	GRU-M-1	LSTM-U-2A	GRU-U-2A
Taux d'apprentissage	0,00001	0,00005	0.00001	0.00001
Taux d'abandon	0,2	0,2	0,2	0
Nombre de couches cachées	512	512	256	256
Taille de la fenêtre d'observation	24 h	24h	48 h	48 h
Nombre d'époques	5	5	5	5
Fonction d'activation	Adam – $Tanh^l$ – sigmoïde	Adam – $Tanh^l$ – sigmoïde	Adam – $Tanh^l$ – sigmoïde	Adam – $Tanh^l$ – sigmoïde

^lTangente hyperbolique

Pour évaluer l'impact des modèles multivariés, et de ce fait l'ajout d'attributs sur la précision de la prédiction du SNR, ils sont comparés aux modèles univariés LSTM-U-2A et GRU-U-2A présentés à la section 3.5.3 (Aladin et al., 2020a).

Ainsi, comme présentées dans le Tableau 4.3, les fonctions d'activation utilisées sont :

- la fonction Adam utilisée comme solveur,
- la fonction tangente hyperbolique (Tanh) utilisée pour déterminer l'état des cellules,
- la fonction sigmoïdale utilisée pour calculer les portes dans chaque bloc LSTM / GRU.

Pour le modèle multivarié LSTM-M-1, le pas d'apprentissage optimal trouvé est 0,00001, tandis que la taille de la couche cachée est 512 neurones, le taux d'abandon 0,2 et la taille de la fenêtre d'observation est 24 heures. Pour le modèle multivarié GRU-M-1, le pas d'apprentissage trouvé est 0,00005, la taille de la couche cachée est 512 neurones, le taux d'abandon 0,2 et la taille de la fenêtre d'observation est 24 heures. Les modèles univariés conservent les mêmes hyper-paramètres trouvés dans l'étude faite à la section 3.4.2.

4.4.2 Évaluation des performances

Une fois entraînés à partir de l'ensemble d'apprentissage et optimisés avec l'ensemble de validation, les modèles ont été évalués dans l'ensemble de test de la base de données KB-1. Cet ensemble de test correspond à une période d'environ 2,5 mois, soit de décembre 2017 à mars 2018. Pour des fins de comparaison, en plus des modèles univariés LSTM-U-2A et GRU-U-2A, un modèle naïf a également été implémenté sur l'ensemble des données de test pour évaluer les performances des modèles multivariés. Ce modèle naïf consistait à attribuer à la valeur de SNR prédite à l'horizon T , la dernière valeur de la fenêtre d'observation N . Tel que soulevé à la section 3.5, d'autres modèles de base peuvent être utilisés pour comparer les performances des modèles. Toutefois, l'objectif principal de cette étude étant de comparer les modèles univariés et les modèles multivariés afin d'évaluer l'impact de l'ajout d'attributs sur

la précision de la prédiction, seul le modèle naïf est utilisé comme base de comparaison. Les métriques utilisées dans l'étape d'évaluation sont :

- le RMSE pour déterminer la précision des modèles,
- l'erreur maximale absolue (*Absolute Maximal Error, AME*) pour évaluer l'impact des mauvaises prédictions sur les performances du modèle,
- le score R^2 pour évaluer la robustesse des modèles,
- le temps de calcul pour évaluer le temps de précision des modèles.

Notons également que ces métriques sont déterminées en exécutant une seule fois les modèles sur l'ensemble de test.

La Figure 4.3 montre les résultats de performance selon ces quatre métriques citées. Notons que l'évaluation de performance des modèles a été réalisée pour des horizons de prédiction allant de 1 h à 96 h.

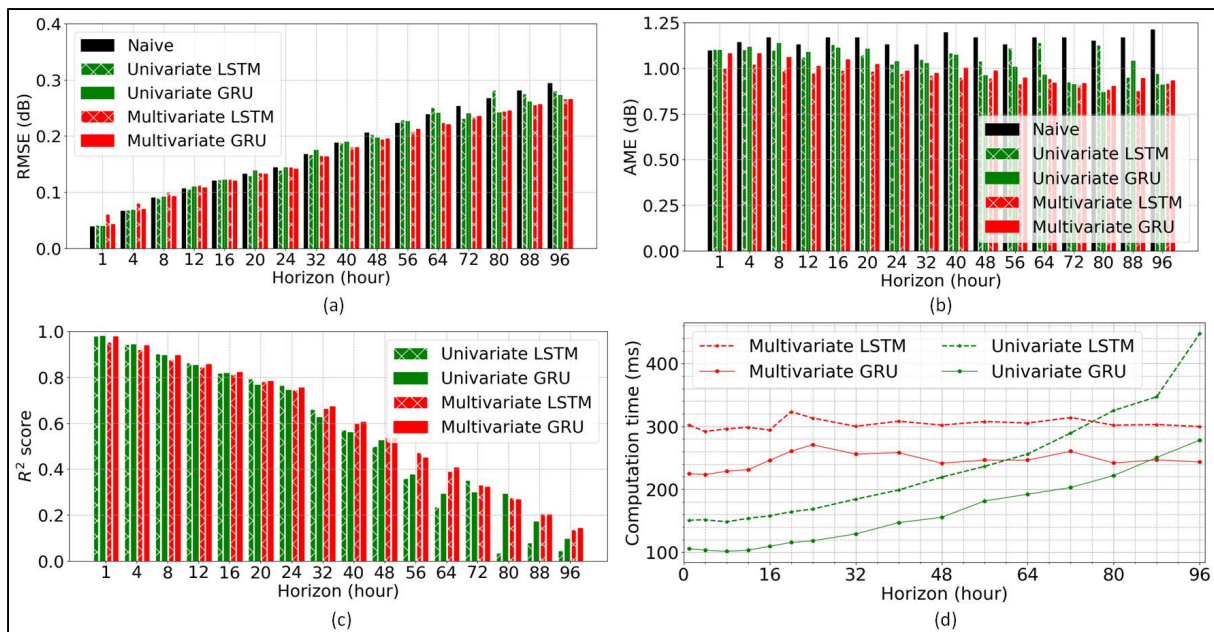


Figure 4.3 Performance des modèles multivariés LSMT-M-1 et GRU-M-1 utilisant l'ensemble de test de KB-1 : (a) RMSE ; (b) AME ; (c) score R^2 ; (d) Temps de calcul
Tirée de Allogba, Aladin et Tremblay (2021)

Le RMSE indique à quel point les valeurs de SNR observées sont proches des valeurs de SNR prédites. Comme le montre la Figure 4.3(a), les valeurs de RMSE augmentent avec l'horizon de prédiction pour tous les modèles. Cela indique que plus l'horizon de prédiction est élevé, plus la performance du modèle est réduite. Par ailleurs, les valeurs de RMSE des modèles sont très similaires pour tous les horizons confondus. Néanmoins, un léger avantage est obtenu par les modèles multivariés sur les modèles univariés et le modèle naïf à mesure que l'horizon de prédiction augmente. En regardant plus en détail, pour les horizons inférieurs à 16 heures, le modèle GRU multivarié surpasse le modèle multivarié LSTM de 0,008 dB à 0,017 dB en RMSE. De même, pour des horizons de prédiction variant de 64 à 96 heures, la différence de RMSE entre les modèles multivariés GRU et LSTM avec la méthode naïve varie respectivement de 0,010 dB à 0,028 dB et de 0,014 dB à 0,027 dB.

La métrique AME est présentée dans la Figure 4.3(b). Les modèles multivariés sont meilleurs que tous les autres modèles. Cela implique que les erreurs de prédiction maximales obtenues à partir des modèles multivariés sont plus petites que celles obtenues à partir des modèles univariés et de la méthode naïve. En effet, les valeurs AME les plus basses ont été atteintes avec les modèles multivariés sur tous les horizons. Celles-ci varient de 0,87 dB à 1,08 dB, par rapport à des valeurs allant de 0,87 dB à 1,13 dB pour leurs homologues univariés. De plus, l'AME de la méthode naïve est la plus élevée sur tous les horizons combinés.

Le score R^2 est présenté sur la Figure 4.3(c). Il décrit dans quelle mesure les caractéristiques sélectionnées expliquent la variabilité des données à prévoir. Comme présentées dans la section 3.5.3, ses valeurs sont inférieures ou égales à 1, de sorte que plus la valeur est proche de 1, plus le modèle est robuste. Le score R^2 confirme les observations issues des courbes du RMSE. Les modèles multivariés GRU-M-1A et LSTM-M-1A fonctionnent de manière très similaire et surpassent généralement les modèles univariés. Pour les horizons inférieurs à 16 heures, les valeurs de R^2 pour les modèles multivariés et univariés sont très proches. Cependant, plus les horizons sont élevés, plus les modèles univariés deviennent instables rapidement avec des valeurs de R^2 moins efficace que celles des modèles multivariés.

Comme le montre la Figure 4.3(d), le temps de calcul est à peu près le même pour les deux modèles multivariés et augmente très légèrement à mesure que l'horizon de prédiction augmente. Parallèlement, pour les modèles univariés, le temps de calcul augmente en fonction de l'horizon de prédiction. De plus, pour des horizons plus longs, les modèles multivariés deviennent plus rapides que les modèles univariés. Les modèles ont été exécutés sur un système doté d'un processeur Intel® Core™ i5-7200U à 2,5 GHz et une mémoire RAM de 8 Go.

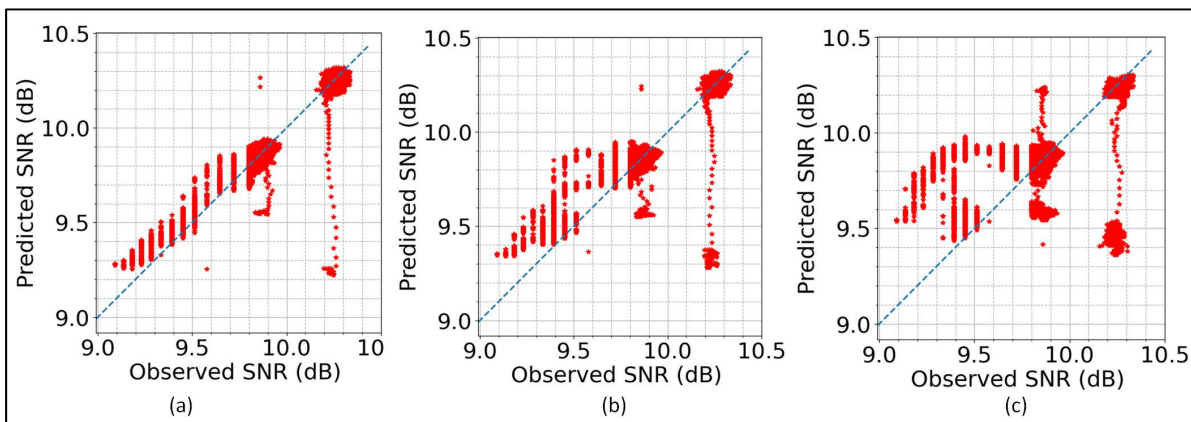


Figure 4.4 SNR prédit vs SNR observé (connexion 1, LSTM-M-1) : (a) horizon de 1 heure ; (b) horizon de 24 heures ; (c) horizon de 96 heures
Tirée de Allogba, Aladin et Tremblay (2021)

La Figure 4.4 montre les données prédites de SNR en fonction des données réelles de SNR pour l'horizon de prédiction la plus courte (1 heure) à l'horizon de prédiction la plus longue (96 heures), ainsi que pour un horizon de prédiction intermédiaire (24 heures) pour le modèle le plus performant, à savoir le modèle LSTM-M-1. Plus les valeurs prédites sont concentrées autour de la ligne pointillée, meilleure est la prédiction. Ainsi, comme le montre la Figure 4.4, les courbes des modèles sont concentrées autour de la ligne pointillée pour les trois horizons présentés. Cela indique le bon fonctionnement des modèles implémentés. Toutefois, des erreurs de prédiction sont observées pour des valeurs de SNR d'environ 10,25 dB. Cela peut être expliqué par le fait que les valeurs de SNR aussi élevées observées dans l'ensemble de test se situent en dehors de la plage de valeurs de SNR présentes dans l'ensemble d'apprentissage, comme le montre la Figure 3.2. Nous pouvons également voir que les courbes sont plus

dispersées plus l'horizon de prédiction augmente, montrant ainsi la diminution des performances des modèles pour les horizons de prédiction plus élevés.

En résumé, les modèles multivariés GRU-M-1 et LSTM-M-1 ont des performances très similaires que leurs homologues univariés pour les horizons inférieurs à 16 heures, avec un léger avantage de performance à mesure que l'horizon de prédiction augmente. Les modèles multivariés sont également meilleurs que le modèle naïf, quel que soit l'horizon de prédiction. Cependant, les erreurs de prédiction AME, plus faibles obtenues avec les modèles multivariés, confirment leur avantage par rapport aux autres modèles (univariés et naïf). De plus, les modèles multivariés sont plus rapides pour prédire les données de SNR que leurs homologues univariés pour les horizons de prédiction à long terme.

4.5 Modèle de prédiction multivarié avec réduction des attributs

Dans cette section, nous étudions l'impact des attributs sur les performances des modèles multivariés dans le but de réduire leur nombre dans l'implémentation des modèles multivariés.

4.5.1 Ingénierie des attributs et présentations des modèles réduits

La première étape pour la réduction des attributs a été d'effectuer une analyse de corrélation sur l'ensemble des attributs de la base de données de KB-1 en utilisant le test de Pearson. Le but de cette analyse est d'évaluer l'importance des différents attributs de KB-1 afin de conserver les plus importantes à utiliser dans les modèles multivariés. Le Tableau 4.4 présente les résultats du test de Pearson.

Tableau 4.4 Analyse de corrélation des attributs de KB-1

	SNR
--	------------

Puissance optique au récepteur	0,8
DGD	-0,2
Température externe	-0,2
Période de la journée	0,01

Le test de corrélation révèle que la puissance du canal est fortement corrélée aux données de SNR (valeur de corrélation = 0,8). Des corrélations beaucoup plus faibles ont été trouvées avec la température, le DGD et la période de la journée.

Sur la base de ces résultats de corrélation trouvés, deux nouveaux modèles multivariés, appelés ici LSTM-M-2 et GRU-M-2, ont par la suite été implémentés en utilisant seulement les données de SNR et de puissance comme attributs. Comme dans la section 4.4.1, les modèles à 2 attributs ont été implémentés à l'aide du package *Keras* dans Python 3. La base de données résultante de la réduction des attributs est divisée en trois groupes selon le ratio (0,64 ; 0,16 et 0,20) correspondant respectivement à l'ensemble d'apprentissage, l'ensemble de validation et l'ensemble de test. À partir du jeu de données de validation, les mêmes hyper-paramètres, à savoir le pas d'apprentissage, le taux d'abandon, la taille de la couche cachée et la taille de la fenêtre d'observation, ont été optimisés en utilisant le RMSE comme critère de choix. Les hyper-paramètres optimaux trouvés pour le modèle LSTM-M-2 sont donc 0,00005 pour le pas d'apprentissage, 512 neurones pour la taille de la couche cachée, 0,2 pour le taux d'abandon et 1 h pour la fenêtre d'observation. Concernant le modèle GRU-M-2, les hyper-paramètres optimaux trouvés sont 0,000025 pour le pas d'apprentissage, 256 pour la taille de la couche cachée, 0,2 pour le taux d'abandon et 8 h pour la taille de la fenêtre d'observation.

4.5.2 Évaluation des performances des modèles réduits

À l'aide de l'ensemble de test, les mêmes métriques de performance (le RMSE, le score R^2 et l'AME), présentées à la Figure 4.5, sont utilisées pour comparer les deux modèles à 2 attributs

LSTM-M-2 et GRU-M-2 au modèle à 5 attributs le plus performant (LSTM-M-1), au meilleur modèle univarié (GRU-U-2A) et au modèle naïf.

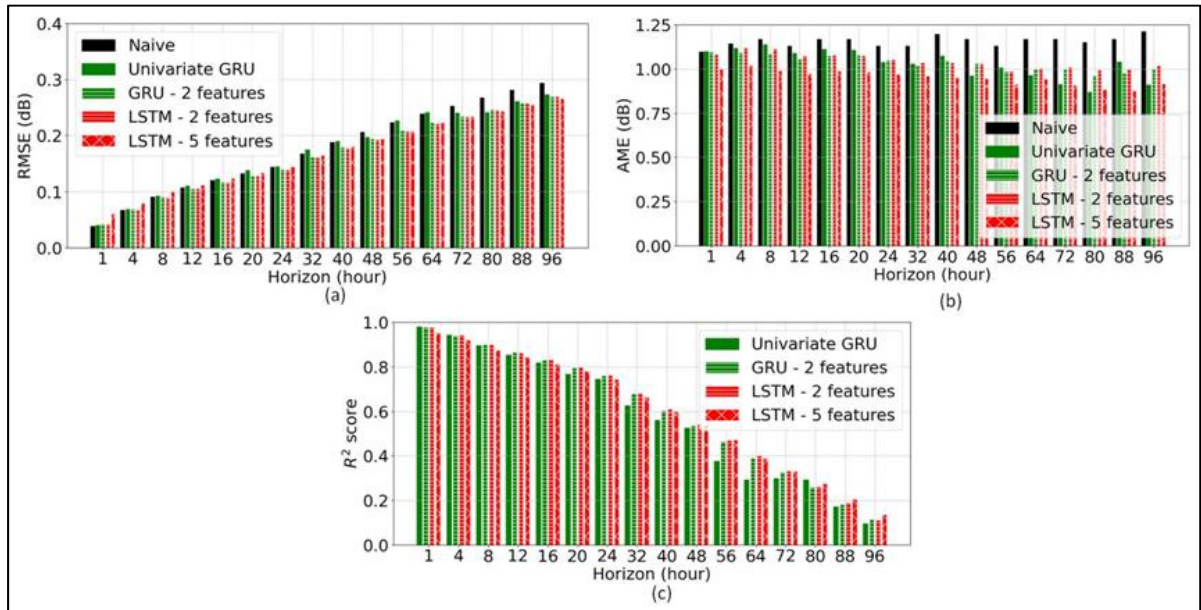


Figure 4.5 Performance des modèles multivariés utilisant l'ensemble de test de KB-1 réduit :
(a) RMSE ; (b) AME ; (c) score R^2

Tirée de Allogba, Aladin et Tremblay (2021); Allogba et Tremblay (2018)

Comme le montrent les courbes de RMSE présentées à la Figure 4.5(a), les modèles LSTM-M-2 et GRU-M-2 sont très similaires au modèle LSTM-M-1. Les deux modèles LSTM-M-2 et GRU-M-2 surpassent par contre le modèle naïf avec une différence RMSE allant jusqu'à 0,023 dB.

La Figure 4.5(b) montre que le modèle LSTM-M-1 surpasse les modèles multivariés à 2 attributs LSTM-M-2 et GRU-M-2 avec une différence en AME d'environ 0,05 dB sur tous les horizons combinés. En revanche, les modèles à 2 attributs performant de manière similaire, mais surpassent le modèle naïf.

La Figure 4.5(c) montre des scores R^2 très similaires pour les modèles multivariés à 2 et 5 attributs.

Les modèles multivariés LSTM-M-2 et GRU-M-2 ont été exécutés sur un système doté d'un processeur Intel® Core™ i5-7200U à 2,5 GHz avec une mémoire RAM de 8 Go. Ils sont restés stables respectivement à environ 15,5 ms et 16,1 ms. En outre, les modèles multivariés LSTM-M-2 et GRU-M-2 sont plus rapides que le modèle LSTM-M-1 à 5 attributs.

Tableau 4.5 Analyse comparative des modèles univariés et multivariés

Métrique	Horizon de prédiction (heures)	Modèle de base	Modèles univariés (Aladin et al., 2020a)		Modèles multivariés			
		Naïf	LSTM	GRU	LSTM-M-1	GRU-M-1	LSTM-M-2	GRU-M-2
		Attributs						
		SNR		SNR, DGD, P_{RX} , Température externe, Période de la journée		SNR, P_{RX} ,		
RMSE (dB)	1	0,04	0,04	0,04	0,06	0,04	0,04	0,04
	24	0,14	0,14	0,14	0,14	0,14	0,14	0,14
	96	0,29	0,28	0,27	0,27	0,27	0,27	0,27
AME (dB)	1	1,10	1,11	1,10	1,00	1,08	1,08	1,10
	24	1,02	1,04	0,97	0,99	1,05	1,05	1,02
	96	0,97	0,91	0,92	0,93	1,02	1,00	0,97

Le Tableau 4.5 présente une analyse comparative des performances des modèles multivariés par rapport aux modèles univariés et au modèle naïf. Les valeurs RMSE et AME sont affichées pour des horizons de prévision de 1 heure, 24 heures et 96 heures, avec les meilleures performances (valeurs les plus basses) indiquées en caractères gras.

La première observation est que les modèles RNN complexes n'ont pas présenté d'avantage RMSE par rapport à une méthode naïve simple pour des horizons très courts (24 heures ou moins). La performance RMSE, légèrement meilleure des modèles univariés par rapport au modèle naïf à l'horizon de 96 heures, ne s'est pas améliorée davantage en utilisant des modèles multivariés.

Cependant, selon la métrique AME, les modèles multivariés à 5 attributs ont surpassé leurs homologues univariés ainsi que le modèle naïf sur tous les horizons de prévision, sauf à l'horizon de 96 heures où le modèle GRU univarié montre un avantage de performance de 0,01 dB. De plus, plus l'horizon de prévision est élevé, plus l'avantage de l'AME est important par rapport aux modèles univariés et au modèle naïf. La meilleure performance globale sur tous les horizons de prévision, en termes de RMSE et d'AME, a été obtenue en utilisant le modèle LSTM multivarié à 5 attributs.

Le modèle GRU multivarié ne doit pas être écarté car il présente une performance similaire (bien que légèrement inférieure) mais un temps de calcul plus court.

Enfin, bien que les modèles multivariés à 2 attributs aient sous-performé par rapport aux modèles multivariés à 5 attributs et aux modèles univariés, ces modèles ne doivent pas non plus être écartés. Les modèles à 2 attributs ont surpassé la méthode naïve et ont présenté un temps d'exécution beaucoup plus court que les modèles multivariés à 5 attributs et les modèles univariés, avec une réduction jusqu'à 95 % du temps de calcul dans les deux cas.

4.6 Généralisation des modèles de prédiction utilisant l'apprentissage par transfert

L'apprentissage par transfert peut être utile lorsqu'il existe une quantité limitée de données pour entraîner les modèles de prédiction. Dans cette section, le concept d'apprentissage par transfert est exploré en testant les modèles univariés et multivariés sous deux scénarios,

utilisant les bases de données KB-2 et KB-3 issues des deux connexions optiques du réseau de CANARIE :

- scénario 1 : le scénario 1 considère la base données KB-2 provenant de la connexion lightpath-2. Celle-ci est transportée dans la même fibre optique que la base de données KB-1 et est déployée sur une section de 600 km de la même route. Rappelons que les métriques de performance utilisées ont été échantillonnées aux 15 minutes sur une période de 5 mois (de février 2017 à juillet 2017). Dans cette étude, la base de données du scénario 1 a été identifiée par KB-2. De même, la connexion cible du scénario 1 est identifiée par lightpath-2 ;
- scénario 2 : la base de données du scénario 2 est constituée des métriques de performance été collectées sur la deuxième connexion optique lightpath-3 pendant la même période de 5 mois (de février 2017 à juillet 2017). Celle-ci est transportée dans la direction opposée de la même route que la connexion lightpath-1. La base de données du scénario 2 sera identifiée par KB-3 et la connexion optique cible du scénario 2 par lightpath-3.

4.6.1 Description de l'approche d'apprentissage par transfert

La Figure 4.6 décrit le concept de l'apprentissage par transfert.

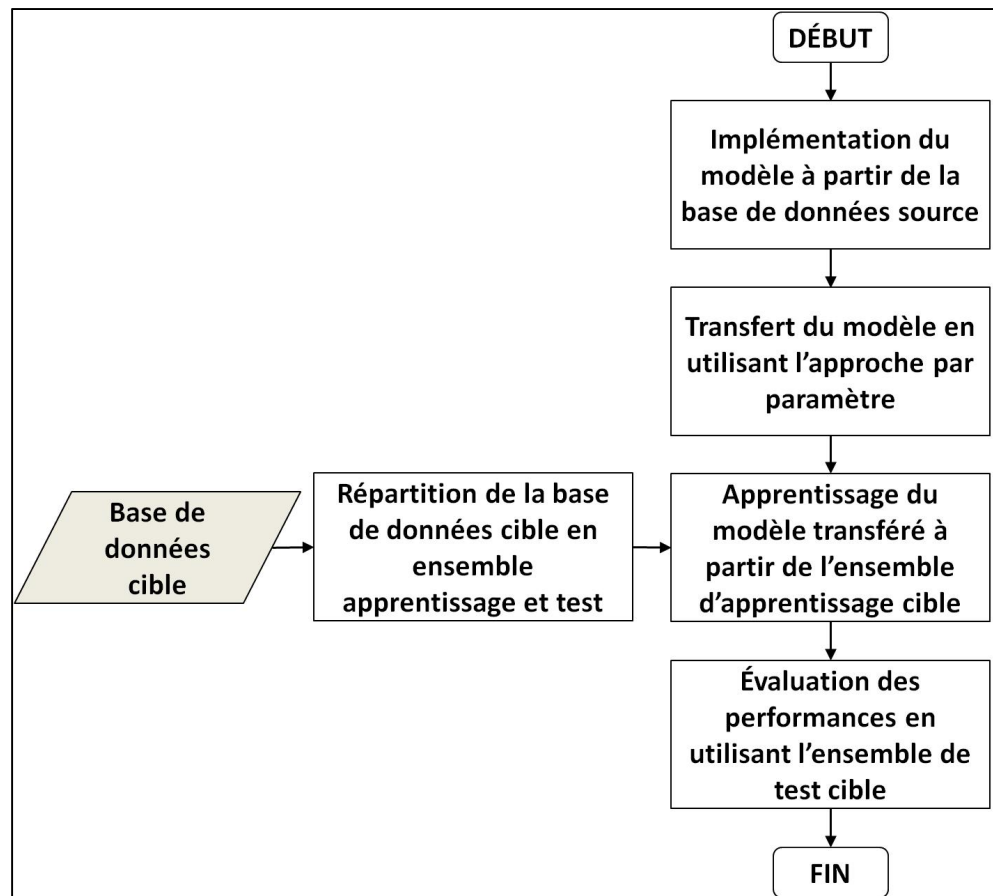


Figure 4.6 Présentation du concept d'apprentissage par transfert

Comme présenté dans la Figure 4.6, le concept d'apprentissage par transfert consiste à transférer un modèle implémenté dans un domaine donné vers un autre domaine pour répondre à un problème prédéfini. Dans la littérature sur l'apprentissage par transfert, un domaine est défini par un espace d'attributs et une distribution de probabilité marginale. Pour un domaine donné, une tâche est définie par un espace d'étiquettes et une fonction prédictive (Weiss, Khoshgoftaar et Wang, 2016). Le domaine dans lequel est implémenté le modèle est appelé domaine source. Parallèlement, le domaine dans lequel le problème doit être résolu est appelé domaine cible. Nous supposons donc que les domaines source et cible sont différents.

En utilisant l'apprentissage par transfert, l'objectif est de déterminer si les modèles prédiction de SNR entraînés avec les données historiques d'une connexion optique spécifique (correspondant au domaine source) peuvent être généralisés et donc utilisés pour prédire le SNR d'une autre connexion optique (correspondant au domaine cible) déployée soit sur une portion de la même fibre optique que le domaine source (scénario 1), soit sur une fibre optique différente du domaine source, c'est-à-dire sur la même route de direction opposée à la connexion du domaine source (scénario 2).

Dans le contexte des présents travaux, les éléments du domaine source sont les éléments de la base de données KB-1 (connexion lightpath-1). Par ailleurs, dans les deux scénarios utilisés pour évaluer le concept d'apprentissage par transfert, les éléments du domaine cible sont issus de deux connexions optiques du réseau de CANARIE, à savoir lightpath-2 et lightpath-3. Comme présenté dans la Figure 4.1, lightpath-3 est transporté entre les mêmes nœuds que lightpath-1, mais dans la direction opposée (scénario 2), tandis que lightpath-2 est transporté dans la même fibre optique sur une section plus courte (600 km) de la même route que lightpath-1 (scénario 1). Notons également, comme vu à la Figure 4.2, que les bases de données cibles ont également des distributions et des variabilités différentes de la base de données source.

4.6.2 Étapes d'implémentation

La première étape d'implémentation des modèles utilisant l'approche d'apprentissage par transfert est le transfert du modèle source vers le domaine cible, comme présenté à la Figure 4.6. Cette étape consiste à utiliser un modèle de prédiction du SNR formé sur une base de données spécifique (dans notre cas KB-1) pour résoudre un nouveau problème. Dans notre cas le nouveau problème est la prédiction des données de SNR des bases de données KB-2 et KB-3. Les modèles de prédiction utilisés dans la démonstration d'apprentissage par transfert sont les modèles les plus performants implémentés précédemment, à savoir le modèle univarié GRU-U-2A et le modèle multivarié LSTM-M-2. Le domaine source représente ici la base de

données KB-1, alors que les deux bases de données KB-2 et KB-3 sont utilisées pour les domaines cibles.

Nous adoptons également l'approche par paramètres. Cette approche suppose que les structures apprises par nos deux modèles entraînés à partir du domaine source peuvent être transférées vers les modèles des deux autres domaines cibles. Ce transfert de structure est effectué à partir d'un partage de poids entre les domaines source (lightpath-1) et cible (lightpath-2 et lightpath-3). L'objectif ici est de réaliser une expérience préliminaire de l'approche d'apprentissage par transfert pour prédire les données de SNR de deux nouvelles connexions optiques issues de deux liaisons optiques différentes en utilisant un modèle préentraîné sur une connexion optique différente. Par conséquent, les poids ont été choisis sans aucune optimisation sur le modèle. Le choix des poids pourrait modifier les performances du modèle, car ces derniers permettent de déterminer les interactions entre tous les neurones de chaque couche et les fonctions d'activation. Ils permettent également de définir les connaissances acquises à partir du domaine source (Mo et al., 2018b).

En outre, en se basant sur la stratégie par ajustement (*fine tune*) présentée par Zhuang et al. (2021), les modèles ont d'abord été implémentés dans le domaine source. Ils ont ensuite été partiellement entraînés dans le domaine cible, comme le montre la Figure 4.6. Ainsi, pour l'implémentation des modèles avec l'approche de l'apprentissage par transfert, chacune des bases de données cible (KB-2 et KB-3) a été divisée selon le ratio 80/20, correspondant à 80% pour la phase d'apprentissage et 20% pour l'évaluation des performances des modèles.

4.6.3 Évaluation des performances

Les métriques utilisées dans la phase d'évaluation sont le RMSE et l'AME. Elles sont présentées à la Figure 4.7.

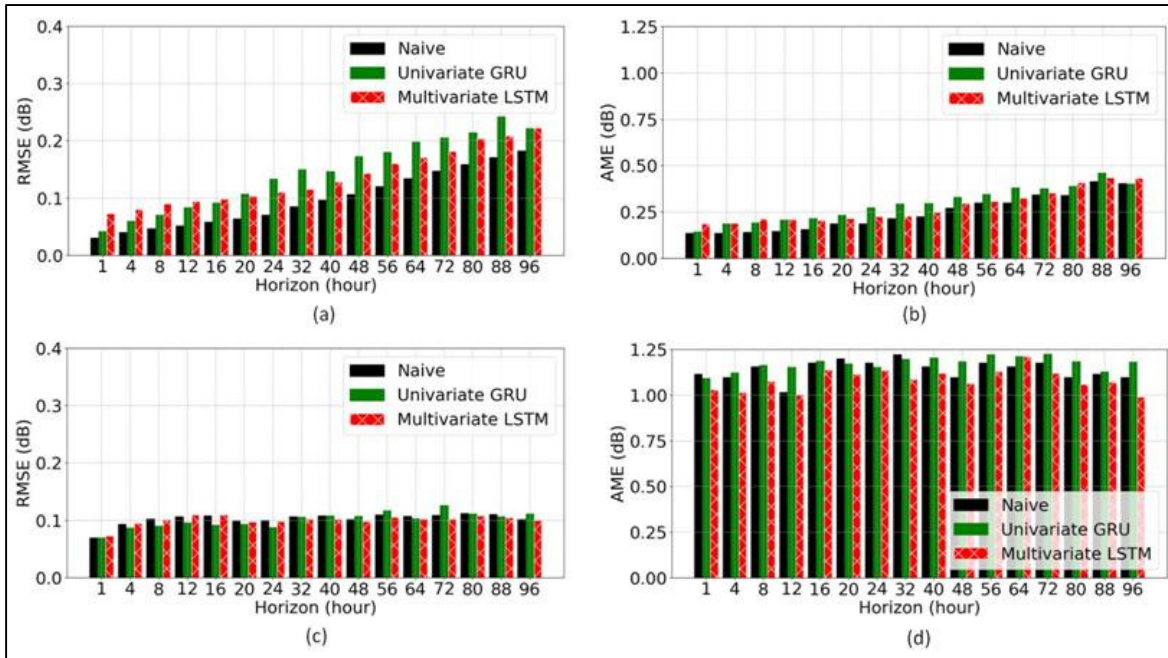


Figure 4.7 Performance des modèles LSTM-M-2 et GRU-U-2A pour la base de données KB-3 (connexion optique de direction opposée sur la même route) : (a) RMSE ; (b) AME – Performances des modèles pour la base de données KB-2 (connexion optique utilisant la même fibre optique sur une portion de la même route) : (c) RMSE ; (d) AME
Tirée de Allogba, Aladin et Tremblay (2021)

Comme le montrent les courbes de RMSE de la Figure 4.7(a), les modèles multivariés et univariés utilisant la base de données KB-3 sous-performent le modèle naïf, tout horizon confondu. Rappelons que le domaine cible lightpath-3 et le domaine source lightpath-1 sont transportés dans des fibres optiques différentes et que leurs bases de données respectives ont des distributions et des variabilités différentes. De plus, dans la et Figure 4.7(c), les modèles multivariés et univariés utilisant la base de données KB-2 ont des performances similaires au modèle naïf pour tous les horizons confondus. Par conséquent, du point de vue du RMSE, les modèles univarié GRU et multivarié LSTM basés sur l'apprentissage par transfert n'ont fourni aucun avantage par rapport au modèle naïf.

Des observations différentes ont été faites en regardant les courbes AME présentées aux Figure 4.7(b) et Figure 4.7(d). En effet, pour la connexion lightpath-3, les modèles multivarié LSTM et univarié GRU basés sur l'apprentissage par transfert ne surpassent pas la méthode naïve. Cependant, pour la connexion lightpath-2, le modèle multivarié LSTM basé sur l'apprentissage par transfert donne de meilleures performances que le modèle naïf, pour tous les horizons confondus. L'avantage AME par rapport au modèle naïf pour l'horizon de prédiction de 96 heures est de 0,11 dB par rapport à 0,304 dB pour le modèle utilisé dans le domaine source lightpath-1. Il est intéressant de noter que les valeurs AME pour les modèles utilisant la connexion lightpath-2 sont du même ordre de grandeur que les modèles utilisant la connexion lightpath-1 (entre 1,00 dB et 1,24 dB). Ce qui n'est pas le cas pour les modèles utilisant la connexion lightpath-3 pour lesquelles les valeurs AME sont nettement inférieures (entre 0,20 dB et 0,49 dB), bien que la largeur de la probabilité de distribution du SNR pour lightpath-2 soit plus étroite que celle pour lightpath-1 et lightpath-3. Cela pourrait potentiellement être expliqué par l'étendue des valeurs de SNR de KB-2 et KB-1 très similaire (respectivement 2,38 dB et 2,35 dB). Ces étendues sont presque deux fois supérieures à l'étendue des valeurs de SNR de KB-3 (1,24 dB).

Enfin les modèles utilisant l'apprentissage par transfert sont plus rapides à exécuter que les modèles classiques pour les deux bases de données, avec des temps de prédiction allant jusqu'à 53 ms pour les modèles utilisant l'approche d'apprentissage par transfert. En outre, les temps d'apprentissage moyens utilisant l'approche d'apprentissage par transfert vont jusqu'à 25 minutes, comparés à 75 minutes pour les modèles utilisant l'approche d'apprentissage de base. Les modèles sont exécutés sur un système avec un processeur Intel® Core™ i5-7200U 2,5 GHz avec une mémoire RAM de 8 Go.

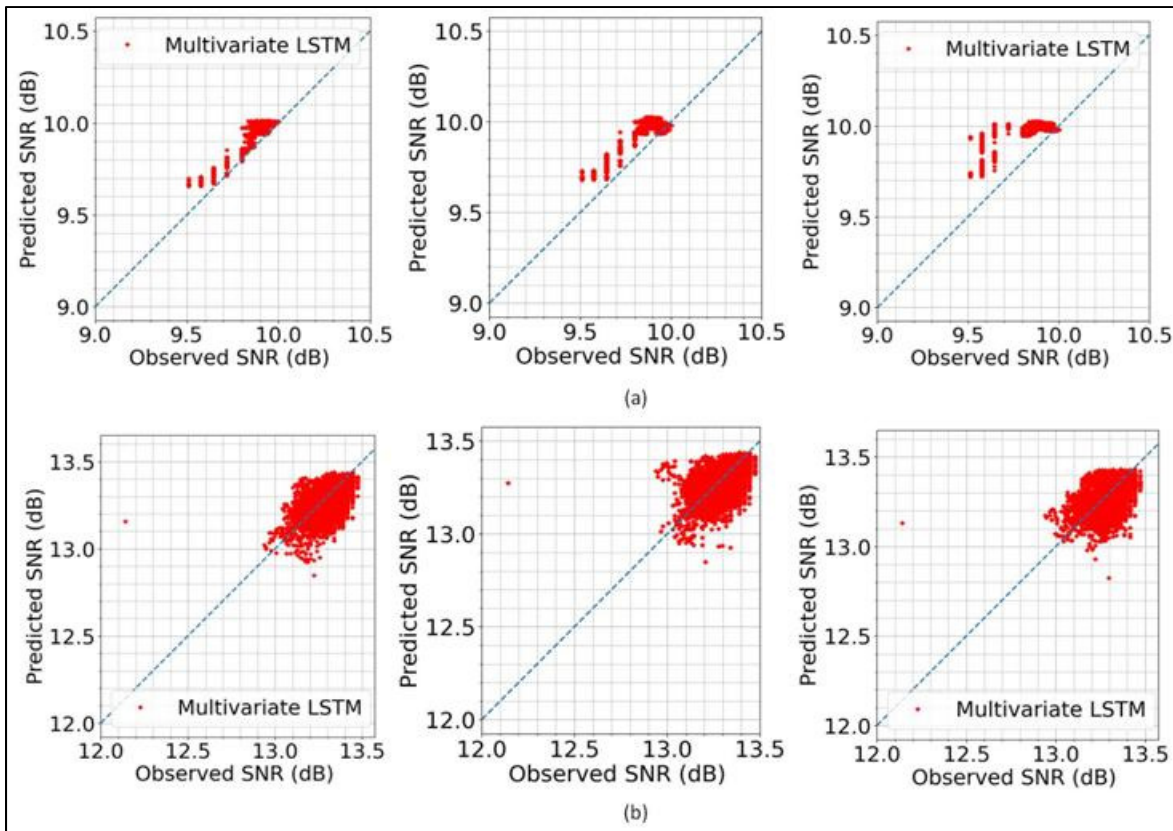


Figure 4.8 SNR prédit vs SNR observé pour l’horizon de 1 heure, 4 heures et 96 heures utilisant le modèle multivarié LSTM basé sur l’apprentissage par transfert : (a) connexion optique cible lightpath-3 (base de données KB-3) ; (b) connexion optique cible lightpath-2 (base de données KB-2)

Tirée de Allogba, Aladin et Tremblay (2021)

La Figure 4.8 montre le SNR prédit en fonction du SNR réel pour l’horizon de prédiction de 1 heure, 4 heures et 96 heures pour le modèle multivarié LSTM basé sur l’apprentissage par transfert. Les courbes des modèles utilisant la base de données KB-2 sont plus concentrées autour de la ligne pointillée que celles utilisant la base de données KB-3. Cela montre une meilleure précision dans la prédiction des modèles utilisant un canal transporté dans la même fibre optique. Notons que le choix des données cibles a un impact sur les performances du modèle. Plus celles-ci ont la même distribution et les mêmes informations que les données source, meilleures seront les performances du modèle. En effet, en choisissant l’approche de transfert des poids des réseaux du domaine source, les domaines cibles héritent des

informations du domaine source. Cette stratégie produit donc des performances satisfaisantes dans les scénarios où il existe une similitude entre les domaines source et cible. Dans notre cas, les éléments de la base de données cible KB-2 proviennent de la même fibre optique que ceux du domaine source KB-1. De même, l'étendue des variables des deux domaines est approximativement la même.

Dans l'ensemble, ces résultats préliminaires montrent le potentiel de réutilisation des modèles de prédiction préalablement formés avec les données d'un canal d'une connexion optique pour prédire les performances d'autres canaux. Les performances du modèle multivarié LSTM-M-2 sont similaires à celles de la méthode naïve, avec un avantage en termes de performance AME.

Le modèle multivarié LSTM semble être de ce fait un candidat prometteur pour prédire le SNR de connexions optiques provenant de la même fibre optique. En effet, même si les performances RMSE sont très proches de la méthode naïve, le modèle multivarié LSTM a fourni de plus petits résultats AME.

4.7 Conclusion

Dans ce chapitre, nous avons proposé deux modèles multivariés LSTM-M-1 et GRU-M-1 pour prédire le SNR d'une connexion optique (lightpath-1) d'un réseau de production en utilisant des métriques de performance collectées sur une période de 13 mois. Ces métriques de performance sont le pré-FEC BER, la puissance reçue du canal, le DGD, la température et la période de la journée. Les modèles LSTM-M-1 et GRU-M-1 ont été implémentés, c'est-à-dire entraînés et testés, en utilisant les 5 attributs composés des métriques de performances de la connexion optique lightpath-1. Les résultats montrent que les modèles multivariés ont obtenu des résultats légèrement meilleurs que leurs homologues univariés et que la méthode naïve, avec des valeurs RMSE inférieures des valeurs de R^2 plus élevées. Par ailleurs, même si les modèles multivariés sont très proches (en termes de RMSE), le modèle multivarié LSTM-M-

1A se distingue pour tous les horizons combinés avec les plus petites valeurs d'erreurs maximales AME avec des écarts allant jusqu'à 0,08 dB, 0,17 dB et 0,28 dB par rapport respectivement aux modèles GRU-M-1, univarié GRU-U-2A et univarié LSTM-U-2A.

En outre, l'analyse statistique effectuée sur l'ensemble des données pour la réduction des attributs a révélé que la puissance du canal était le meilleur attribut additionnel (parmi les attributs disponibles) à utiliser avec les valeurs passées de SNR pour prédire le SNR sur des horizons allant jusqu'à 96 heures.

Enfin, le meilleur modèle multivarié construit à l'aide du KB-1, à savoir le modèle multivarié LSTM-M-2, a été réentraîné sur de plus petits ensembles de données (KB-2 et KB-3) provenant de deux connexions optiques différentes lightpath-2 et lightpath-3, correspondant respectivement au scénario 1 et scénario 2. Dans le scénario 1, la connexion lightpath-2 provenait d'une plus petite portion de la même route que la connexion lightpath-1 (dans laquelle est issue KB-1) alors que, dans le scénario 2, la connexion optique lightpath-3 provenait d'une connexion optique transportée dans une fibre optique différente et une direction opposée à la route utilisée par lightpath-1 (connexion B présentée à la section 3.3.1). Cette approche a été réalisée afin d'explorer le concept de l'apprentissage par transfert. Avec des performances RMSE similaires au modèle naïf, il a ainsi été démontré qu'il est possible de prédire le SNR de la connexion lightpath-2 transportée dans la même fibre optique que la connexion du domaine source (lightpath-1), avec un avantage AME allant jusqu'à 0,13 dB. Cependant, cette même assertion n'a pas pu être validée dans le cas de la prédiction du SNR de la connexion lightpath-3 déployée sur la même route, mais de direction opposée à la connexion du domaine source (lightpath-1), avec une différence RMSE avec la méthode naïve allant jusqu'à 0,051 dB, équivalent à une baisse de prédiction d'environ 23%. Avec de meilleures performances AME obtenues, l'apprentissage par transfert apparaît donc comme une solution prometteuse pour les horizons de prédiction à court terme avec nos modèles multivariés par rapport au modèle naïf. En outre, en comparant au modèle d'apprentissage

classique, l'apprentissage par transfert permet de réduire le temps d'apprentissage des modèles d'environ 6 fois et le temps de calcul des modèles d'environ 3 fois.

Dans l'ensemble, notre modèle multivarié LSTM-M-2 a mieux performé que le modèle multivarié GRU et les modèles univariés LSTM-U-2A et GRU-U-2A pour la prédiction du SNR d'une connexion optique pour des horizons de prédiction allant jusqu'à 96 heures. De plus, le temps de calcul des modèles multivariés n'augmente pas en fonction de l'horizon, contrairement au temps de calcul des modèles univariés qui augmentent avec l'horizon de prédiction.

Les algorithmes d'apprentissage machine peuvent également être utilisé pour la détection anticipée des anomalies observées dans les performances des connexions optiques. Cette détection anticipée permettrait aux opérateurs de déployer des solutions avant que la connexion ne se dégrade. Ainsi, le prochain chapitre présente nos travaux portant sur l'implémentation d'une méthode de détection anticipée des anomalies de BER observées dans les connexions optiques.

CHAPITRE 5

MODÈLE DE DÉTECTION ANTICIPÉE DES ANOMALIES DE PERFORMANCE DE CONNEXIONS OPTIQUES

5.1 Introduction

Les liaisons optiques déployées aujourd'hui transportant un trafic de plus en plus important pouvant atteindre plusieurs Térabits par secondes par fibre optique, les pannes y survenant entraînent d'immenses pertes de données. Celles-ci peuvent être classées en deux groupes : les pannes matérielles (vieillissement de la fibre optique, interruptions de service, etc.) et les pannes logicielles (dysfonctionnement du système, détérioration des performances des équipements, etc.). Ces pannes sont également caractérisées par une dégradation des métriques de performance du réseau (Chen et al., 2019). Afin de limiter l'impact de ces pannes dans le réseau, les principales solutions s'orientent sur l'analyse des métriques de performance collectées dans le réseau. Celles-ci permettent en effet à la fois de diagnostiquer les causes des problèmes survenant dans le réseau ou de déployer des équipes de maintenance après une panne.

La première solution est donc de prédéfinir des marges dans les systèmes dans le but de prévenir une dégradation dans les métriques de performance du réseau. Cependant, ces marges peuvent être trop élevées, entraînant une utilisation excessive des ressources du réseau, ou trop faibles, rendant la probabilité de panne plus élevée et donc une présence d'alarme importante (Chen et al., 2019; Pointurier, 2017; Wang et al., 2017). La seconde solution afin d'éviter les pannes est l'utilisation des techniques d'apprentissage machine pour la prédiction ou la détection des pannes logicielles ou matérielles (Francesco Musumeci et al., 2018; Gu, Yang et Ji, 2020). Notons que la prédiction des pannes vise à analyser et à exploiter les données historiques collectées dans le réseau afin de prédire une défaillance avant la dégradation des performances ou des équipements du réseau. La détection des pannes, par ailleurs, vise à détecter ou identifier la dégradation des équipements ou des performances du réseau à partir

d'algorithme de classification (Zhang et al., 2020). Ces deux méthodes ont pour but de permettre aux opérateurs de déployer des solutions proactives à temps afin d'éviter la panne ou la défaillance du système et garantir ainsi une meilleure fiabilité du réseau.

Nous proposons, dans ce chapitre, un modèle de détection des pannes logicielles basé sur la détection des anomalies observées dans les données de SNR collectées sur connexions optiques existantes d'un réseau optique de production d'un opérateur majeur aux États-Unis. Le modèle est constitué de deux modules, à savoir un module de définition des anomalies et un module de détection des anomalies. Le premier module permet la définition et l'extraction des attributs et des étiquettes à partir de l'approche de l'étendue interquartile (*interquartile range*, IQR). Le second module est un modèle de détection des anomalies basé sur l'algorithme SVM. L'algorithme SVM se présente ici comme l'un des algorithmes les plus performants en termes de précision, de complexité et de rapidité dans le domaine de la détection d'anomalies (Chandola, Banerjee et Kumar, 2009; Musumeci et al., 2019a).

Les résultats présentés dans ce chapitre font l'objet d'un article de revue en cours de préparation: S. Allogba, B. L. Yameogo and C. Tremblay, « Extraction and Early Detection of Anomalies using Machine Learning Models », to Journal of Lightwave Technology (version révisée soumise le 20 octobre 2021).

Le chapitre est organisé comme suit. La section 5.2 présente une revue des travaux effectués dans le cadre de la détection des pannes logicielles et matérielles. On y trouve également l'objectif de ce travail ainsi que les étapes d'implémentation du modèle de détection des anomalies. La section 5.3 présente le module de définition des attributs et des étiquettes. Elle inclut la présentation de la base de données, les étapes de prétraitement et de présentation des attributs, et la présentation de la notion d'anomalies dans le cadre de notre étude. La section 5.4.1 couvre l'étape de l'ingénierie des attributs. Elle présente également les méthodes pour sélectionner les meilleurs attributs pour utiliser dans le module de détection des anomalies de

SNR. Les détails d'implémentation du module de détection des anomalies, ainsi que les résultats de performance, sont présentés à la section 5.4.2.

5.2 Définition du problème

5.2.1 Présentation des travaux

Dans le cadre de nos travaux, nous nous concentrerons sur les problèmes de détection des pannes matérielles et logicielles. Wang et al. (2017) proposent un modèle outil de détection des pannes matérielles basé sur l'algorithme. Cet outil a pour but de classifier la présence d'une panne ou non dans le réseau en observant l'état des attributs caractérisant les pannes. Notons que les attributs utilisés sont définis par les paramètres physiques de l'équipement (puissance optique d'entrée, puissance optique de sortie, courant laser, température du laser), les propriétés statistiques du BER (valeurs maximales, valeurs minimales et valeurs moyennes), et la température environnementale. En outre, leur état est prédit sur un horizon d'un jour en utilisant le lissage exponentiel double. Wang et al. (2017) définissent les étiquettes, correspondant aux pannes par un seuil arbitraire correspondant au temps d'inutilisation de l'équipement. Ce seuil arbitraire fixé dépend fortement de l'expert et peut ne pas être représentatif des pannes matérielles observées dans un contexte réel.

Par ailleurs, Vela et al. (2017a; 2017b; 2017c) ont concentré leurs travaux sur la détection des pannes logicielles. Les auteurs ont implémenté un algorithme à état fini et utilisant également les réseaux bayésiens pour la détection d'anomalies de BER caractérisées par les changements d'état inattendus observés dans les données de BER. Les étiquettes caractérisant ces anomalies de BER sont définies par un seuil fixe calculé à partir de la valeur moyenne et de l'écart-type des données de BER observées. Par ailleurs, les données de BER utilisées sont des données synthétiques issues d'un banc de test expérimental d'une connexion de 320 km. De plus, les attributs caractérisant les anomalies de BER sont les paramètres statistiques de ces données sur une période d'observation variant de 2 minutes à 10 jours.

Shahkarami et al. (2018) comparent trois algorithmes différents, à savoir le SVM, le RF et les réseaux de neurones pour détecter les pannes logicielles. Ils définissent également les pannes logicielles par un changement d'état significatif de la variation des données de BER pendant une période d'observation prolongée. Cette période d'observation varie de 18 à 73 minutes. Ce changement d'état est donc considéré comme significatif lorsque celui-ci dépasse un seuil fixé par l'expert. Les résultats montrent que l'algorithme RF surpasse les autres algorithmes. Tout comme les travaux de Vela et al, les auteurs ont testé leurs modèles sur des données synthétiques de BER issues d'un banc de test expérimental d'une connexion de 380 km et les attributs considérés sont les paramètres statistiques de ces données.

Toujours dans le contexte de détection des pannes logicielles dans le réseau, Chen, et al. (2019; 2018) ont proposé un modèle basé sur les réseaux de neurones pour détecter les anomalies observées dans les données de puissance d'une connexion optique. Notons que ces anomalies sont caractérisées par les distorsions observées dans la distribution des valeurs de puissance. Ainsi, les auteurs définissent premièrement les étiquettes qui seront implémentées dans leur modèle en utilisant une méthode de regroupement des données basé sur la densité. Cette méthode, contrairement aux études précédentes, permet aux auteurs de définir des anomalies sans se baser sur un seuil fixe. Par la suite, le réseau de neurones est appliqué en utilisant les densités, fournies par la méthode de regroupement, comme attributs. Par ailleurs, tout comme les travaux précédents, les auteurs ont testé leurs modèles sur des données synthétiques issues d'un banc de test expérimental d'un réseau de 220 km. Seulement 6 connexions de ce réseau conçu expérimentalement ont été utilisées.

Pour résumer, les caractéristiques des différents modèles implémentés pour la détection des pannes dans le réseau présentent trois aspects communs. Le premier aspect est l'utilisation de base de données synthétique pour l'implémentation des modèles. Chacune des données utilisées a été extraite de connexions établies expérimentalement et variant de 200 km à 380 km. Le second aspect se situe dans la définition des anomalies. Celle-ci a été basée sur

l'approche par seuil fixe défini par l'expert. Cependant, du fait du caractère aléatoire des anomalies et de leur occurrence peu fréquente, les deux premiers aspects ne sont pas représentatifs des contextes réels du réseau. Finalement, le troisième aspect réside dans le choix des attributs. Ceux-ci sont constitués principalement des propriétés statistiques des données observées (valeurs minimales, valeurs maximales, valeurs moyennes, etc.) durant une période d'observation de l'ordre de quelques heures. Seuls les travaux de Chen, et al. (2019; 2018) ont des caractéristiques différentes pour les deux derniers aspects.

5.2.2 Objectif

Pour renforcer la performance des modèles de détection des anomalies, la première solution serait d'utiliser des méthodes de définition des anomalies différentes des méthodes par seuil fixe proposées dans la littérature. Ces méthodes seraient basées soit sur des approches d'extraction de similitudes (méthode non supervisée) ou des méthodes à seuil variable utilisant des approches statistiques (méthode supervisée). La seconde solution se trouverait dans le choix et la diversité des attributs. L'objectif de cette étude est de valider s'il est possible de détecter les pannes dans un réseau à partir d'un modèle à deux fonctions. La première fonction permet de définir automatiquement les anomalies observées dans les données de SNR d'une connexion optique. La seconde fonction permet de détecter de façon anticipée les anomalies définies. Cette étude intègre donc les deux solutions précitées permettant d'améliorer la performance des modèles de détection des anomalies. De plus, elle est réalisée en utilisant des données de terrain, permettant ainsi d'obtenir des modèles de détection des anomalies plus réalistes et conformes à la réalité du terrain.

La Figure 5.1 présente la méthodologie utilisée pour l'implémentation du modèle de détection des pannes dans le réseau. Chacune de ces étapes est détaillée dans les sections suivantes. Les étapes 1 et 2 font partie du module de définition des anomalies tandis que les étapes 3 et 4 font partie du module de détection des anomalies :

- analyse et prétraitement des données (étape 1) : l'une des particularités de ce modèle est le fait qu'il utilise comme attributs les éléments issus de la décomposition des séries temporelles, les métriques de performances collectées du réseau et les facteurs externes au réseau (activités humaines). Ainsi, cette étape consiste donc à analyser et nettoyer les données brutes collectées afin d'en extraire les attributs à utiliser pour la construction du modèle ;
- définition et extraction des anomalies (étape 2) : elle a pour but de définir les étiquettes de la base de données (anomalies) sans tenir compte d'un seuil fixé par un expert, mais seulement du comportement des données ;

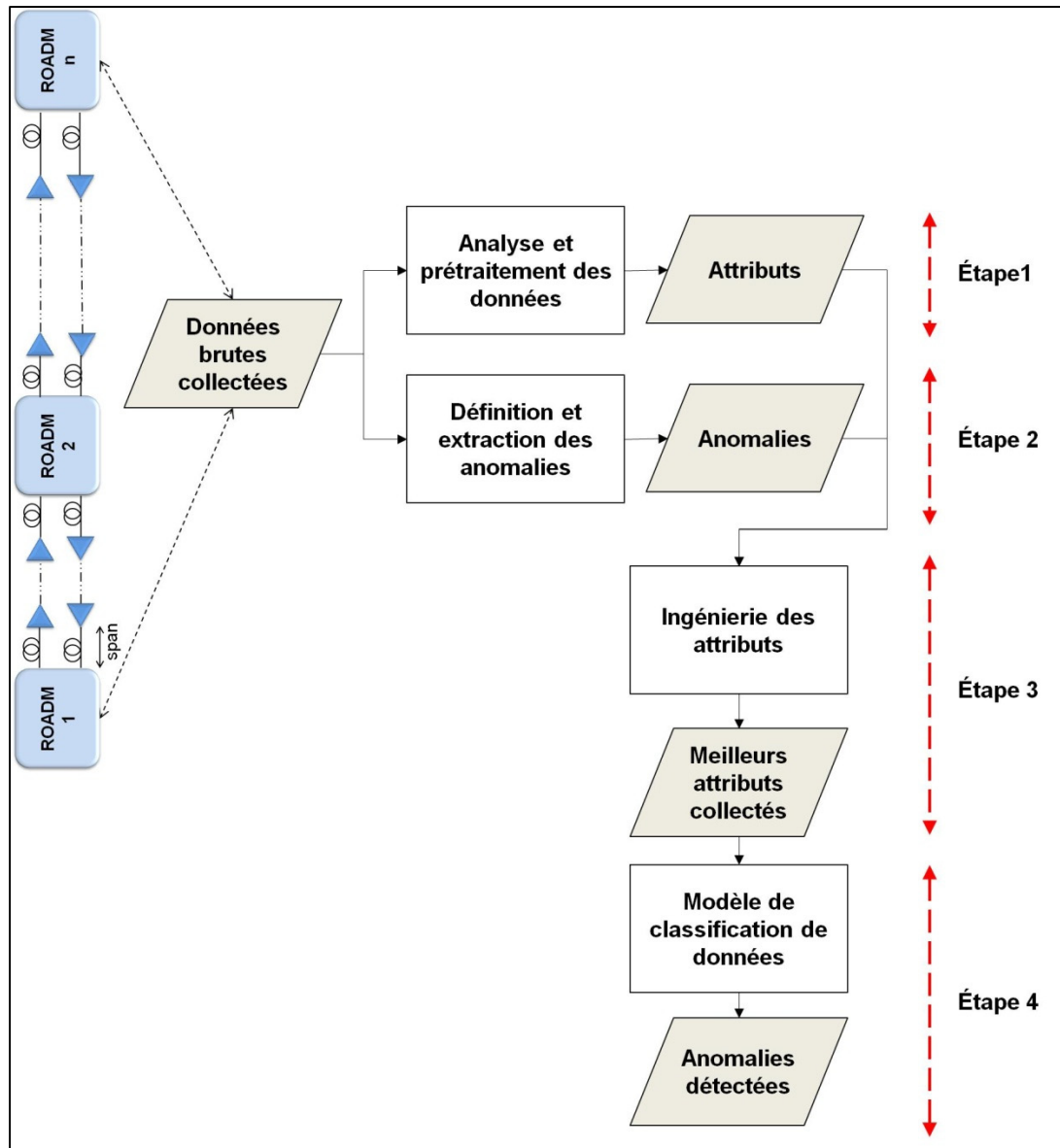


Figure 5.1 Étapes d'implémentation du modèle de détection des anomalies dans les données de SNR de connexions optiques

- ingénierie des attributs (étape 3) : elle a pour but de sélectionner l'attribut ayant le plus d'impact sur les anomalies observées dans le SNR. En effet, rappelons que le choix des attributs a un impact sur les performances des modèles d'apprentissage machine (Zhang et al., 2020). Trois techniques d'apprentissage machine sont utilisées dans cette étape : la

méthode Boruta, la méthode RF et la méthode d'élimination récursive des attributs (*Recursive Feature Elimination, RFE*) ;

- implémentation du modèle (étape 4) : en utilisant l'algorithme SVM, cette étape permet de détecter les anomalies de SNR. Elle inclut également la phase d'optimisation du modèle.

En outre, les attributs trouvés à l'étape 3 sont utilisés lors de la construction du modèle. Ainsi, quatre modèles de détection anticipée d'anomalies SNR sont présentés. Chacun d'entre eux a une combinaison différente d'attributs. Cela a pour but d'évaluer l'impact des attributs sur les anomalies SNR.

5.3 Module de définition des anomalies

Cette section présente premièrement une description du réseau dans lequel sont extraits les éléments constituant la base de données. Par la suite, le module de définition des anomalies est détaillé. Celui-ci comprend la définition des anomalies dans le contexte de notre étude et leur procédure d'extraction, ainsi que l'extraction des attributs.

5.3.1 Description du réseau

La base de données à l'étude est construite à partir des métriques de performance de 8 connexions optiques établies sur une portion du réseau de production d'un opérateur aux États-Unis. Notons que le réseau étudié est constitué de 12 routes optiques à la fois enfouies et aériennes variant de 100 km à 1400 km. Il est composé également de sites d'amplification espacés de 25 à 50 km environ et de sites de ROADM. Dans un souci de confidentialité, les informations sur les sites d'amplification et les sites de ROADM ne sont pas fournies. Toutes les connexions optiques sont des canaux PM-QPSK à 100 Gbit/s.

Les métriques de performance recueillies pour la construction de la base de données sont les données de puissance optique reçue au récepteur (P_{RX}) et les données de SNR au récepteur,

collectées à une fréquence d'échantillonnage de 15 minutes durant une période d'observation de 12 mois.

La base de données finale comprend donc 259 837 éléments de chaque métrique. Les séries temporelles comprennent des données manquantes, à savoir 282 éléments. Ainsi, tout comme dans les chapitres précédents, celles-ci ont été remplacées par une valeur fixe correspondant à la valeur médiane des données observées. Il est tenu pour acquis ici que le fait de choisir une valeur fixe n'affecte pas la distribution des données à cause du faible pourcentage de données manquantes par rapport à la base de données totale (0,11%).

5.3.2 Définition des anomalies

Le module de définition des anomalies est divisé en deux parties. La première partie consiste à définir et identifier les étiquettes à attribuer aux éléments de la base de données. En d'autres termes, cela revient à définir et extraire les anomalies de SNR. Aussi, la définition des anomalies résulte-t-elle de l'exploration des données de SNR.

De manière générale, la notion d'anomalie est fonction du domaine d'application et représente habituellement un ensemble de points de la base de données non conformes au comportement normal du domaine. Le comportement normal, quant à lui, est défini par un intervalle dans lequel se trouve la majorité des valeurs (Chandola, Banerjee et Kumar, 2009).

Dans la présente étude, la définition des anomalies a été faite à partir de l'approche de l'intervalle interquartile (*interquartile range, IQR*), aussi appelée règle de répartition par boîte à moustaches. En effet, cette approche permet de déterminer l'intervalle correspondant au comportement normal, et de ce fait les anomalies observées dans les données. Notons que cette approche a été appliquée dans les systèmes de détection d'intrusion et le domaine informatique (Chandola, Banerjee et Kumar, 2009; Vinutha, Poornima et Sagar, 2018). Elle consiste à estimer des valeurs seuils au-delà desquelles les anomalies sont identifiées en se basant sur des

mesures statistiques de dispersion. Ces valeurs seuils sont identifiées par la plage quartile interne (25% et 75%) déterminée à partir de la distribution autour de la médiane, tel que présenté à la Figure 5.2.

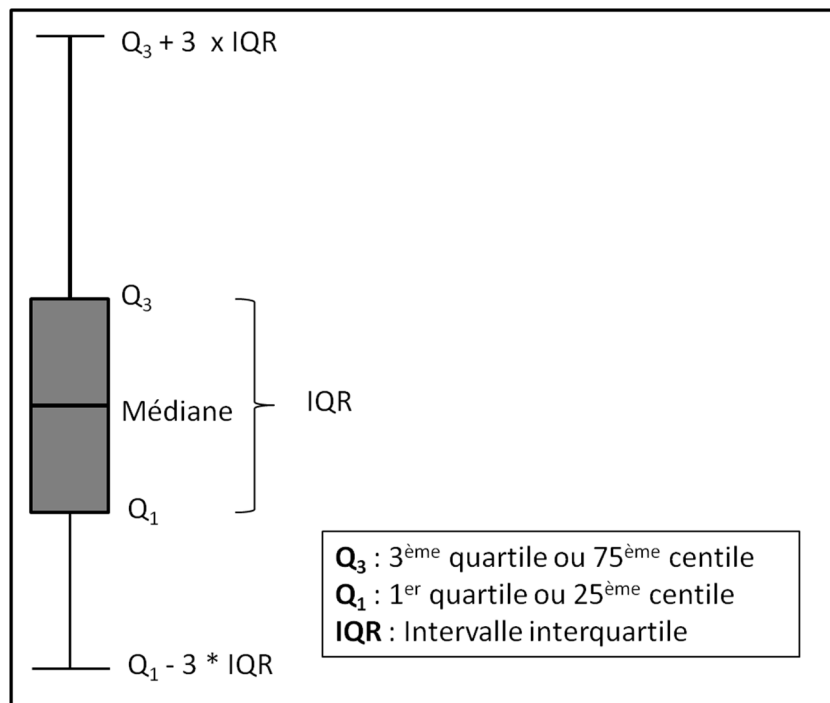


Figure 5.2 Présentation de l'approche IQR pour la détermination des anomalies

En d'autres termes, les anomalies sont identifiées en suivant deux étapes. La première étape consiste à calculer l'intervalle IQR, comme présentée dans l'équation (5.1). Cet intervalle correspond à la différence entre le 75^{ème} centile et le 25^{ème} centile également appelée respectivement le quartile supérieur et le quartile inférieur.

$$IQR = Q_3 - Q_1 \quad (5.1)$$

où

- Q_3 représente le troisième quartile ou le 75^{ème} centile,

- Q_1 représente le premier quartile ou le 25^{ème} centile.

La seconde étape est l'établissement du seuil de définition d'une anomalie. Celui-ci est établi en élargissant la ligne de base 25/75 d'un facteur de 3 IQR. En dessous de cette ligne de base, tel qu'illustré à la Figure 5.2, toute instance de données est identifiée comme anomalie. Cela est traduit par l'équation (5.2) :

$$anomalie \leq Q_1 - 3 * IQR \quad (5.2)$$

Le « package » *anomalize* de Rstudio a été utilisé afin d'implémenter l'approche IQR dans le module de définition des anomalies (Dancho et Vaughan, 2018). La Figure 5.3(a) montre les anomalies obtenues pour les 8 connexions optiques en utilisant l'approche IQR et la Figure 5.3(b) présente un exemple d'anomalies dans les données de SNR pour une des connexions optiques.

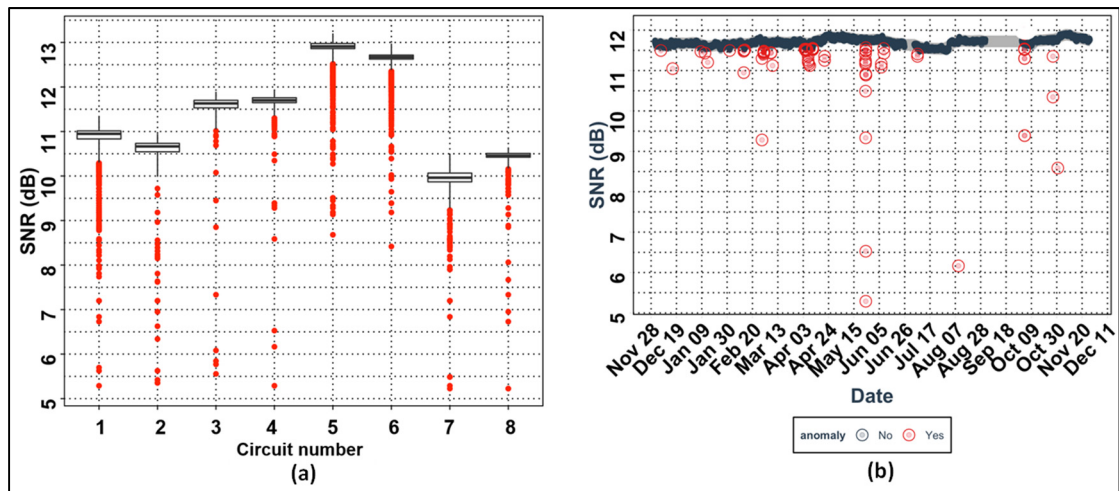


Figure 5.3 (a) Identification des anomalies des 8 connexions optiques en utilisant l'approche IQR ; (b) Exemple d'anomalies observées sur les données de SNR d'une des 8 connexions optiques

Tirée de Allogba, Yaméogo et C. (2021)

Notons que la durée de ces anomalies (ou changements brusques du SNR) varie généralement de 15 minutes à moins de 2 heures.

Tel que présenté à la Figure 5.3, toute valeur située dans la zone inférieure de l'intervalle IQR est considérée comme anomalie. Ainsi, à partir de la définition et de l'extraction des anomalies, deux classes sont définies : la classe 0 (comportement normal ou non-anomalie) et la classe 1 (anomalie). Le module de définition des anomalies permet donc d'extraire de la base de données 257 739 éléments étiquetés « non anomalie » et 2 098 éléments étiquetés « anomalie ». Les éléments étiquetés « anomalie » représentent moins de 1% de la base de données. Cette forte disproportion des données pourrait affecter la précision des modèles en surestimant la précision de détection. Par conséquent, pour faire face au jeu de données déséquilibré, le nombre d'instances « non anomalie » a été réduit (de manière aléatoire) à environ 4 % , comme proposé par Zhang et al. (2020). La base de données finale comprend donc 10 490 éléments étiquetés « non anomalie » et 2 098 éléments étiquetés « anomalie ».

5.3.3 Identification des attributs

La seconde partie du module de définition des anomalies est l'identification et l'extraction des attributs à partir des données extraites du réseau. Ainsi, les attributs issus du module de définition des anomalies sont regroupés en trois catégories.

La première catégorie est issue de la décomposition des séries temporelles (*Seasonal and Trend decomposition, STL*) des données de SNR collectées dans le réseau. Celle-ci est effectuée par régression pondérée locale (*Locally Estimated Scatterplot Smoothing, LOESS*), tel que décrit dans (Yaméogo et al., 2020), et consiste à séparer la série temporelle en trois composantes, représentant les premiers attributs :

- la composante tendance (tendance, T) : elle représente la tendance globale des variations générales de la série temporelle, soit à la hausse, à la baisse ou nulle pendant la période d'observation. Notons que la tendance n'est pas toujours la même pour différentes tranches de la période d'observation. En d'autres termes, pour un découpage donné dans une période d'observation globale, on peut observer une tendance des données à diminuer pendant un

- découpage de cette période alors que pour un autre découpage temporel on peut observer une tendance à augmenter ou à rester stable ;
- la composante saisonnière (saison, S) : elle représente la périodicité observée dans la série temporelle, c'est-à-dire les variations qui se répètent sur une certaine période de temps (un trimestre de l'année, le mois, un jour de la semaine, etc.) ;
 - la composante irrégulière (résidu, R) : elle représente le résidu ou le reste de la série temporelle après extraction des autres composantes. Elle suit une distribution normale avec une moyenne nulle. Cette composante peut être utilisée pour repérer les événements aléatoires observés dans les données (Yaméogo et al., 2020).

La seconde catégorie des attributs est issue des informations temporelles dérivées de la période d'observation. En effet, les effets temporels ont un impact sur le comportement d'une série temporelle. Les informations temporelles retenues sont donc : l'heure(h), le jour de la semaine (J), le numéro de la semaine dans l'année ($ind_{semaine}$), et la période de la journée ($période_{jour}$) caractérisée par le jour, le matin, le soir et la nuit. Notons que l'extraction des informations temporelles a été effectuée à partir du paramètre *datetime* dans Rstudio avec la librairie *lubridate*.

Finalement, la dernière catégorie des attributs est issue des métriques de performance collectées dans le réseau, à savoir la puissance P_{RX} et le SNR. Les derniers attributs ajoutés sont donc les valeurs minimales et maximales du SNR (notées respectivement min_{SNR} et max_{SNR}) ainsi que la puissance optique au récepteur du canal.

5.4 Module de détection des anomalies

5.4.1 Ingénierie des attributs

Le module de définition des anomalies a permis d'identifier 10 attributs à savoir : la saison (S), la tendance (T), le résidu (R), les valeurs minimum (min_{SNR}) et maximum (max_{SNR}) du SNR,

la puissance optique au récepteur (P_{RX}), le jour (J), le numéro de la semaine dans l'année ($ind_{semaine}$), l'heure (h) et la période de la journée ($période_{jour}$). Notons que les attributs peuvent ajouter soit des informations pertinentes dans la description des étiquettes, soit du bruit, ou des informations redondantes dans la description des étiquettes (Chandrashekar et Sahin, 2014). Ainsi, l'ingénierie des attributs comprend l'analyse et la sélection des meilleurs attributs à utiliser dans le modèle pour la détection des anomalies. Elle permet de plus de réduire le temps d'exécution du modèle finale et d'améliorer ses performances.

5.4.1.1 Phase de nettoyage des attributs

La première étape issue de l'ingénierie des attributs consiste à éliminer les attributs dépendants entre eux. Pour ce faire, une analyse de corrélation a été effectuée en utilisant le test de Pearson. Cette analyse a pour but de mettre en évidence les attributs fortement corrélés entre eux. La Figure 5.4 présente les résultats de l'analyse de corrélation.

S	1									
T	0,2	1								
R	-0,3	-0,2	1							
Min _{SNR}	0,9	0,6	-0,3	1						
Max _{SNR}	0,9	0,6	-0,3	1	1					
P _{RX}	0,01	0,05	0,07	0,03	0,03	1				
J	-0,005	-0,01	0,002	-0,009	-0,009	-0,005	1			
h	0,001	-0,006	-0,009	-0,002	-0,002	-0,006	-0,01	1		
ind _{semaine}	0,1	0,5	-0,2	0,3	0,3	0,04	0,002	-0,0002	1	
période _{jour}	-0,001	-0,02	-0,007	-0,008	-0,008	0,001	-0,02	0,3	-0,009	1
	S	T	R	Min _{SNR}	Max _{SNR}	P _{RX}	J	h	ind _{semaine}	période _{jour}

Figure 5.4 Résultats de l'analyse de corrélation des attributs pour le modèle de détection d'anomalies

Adaptée de Allogba, Yaméogo et C. (2021)

Comme on peut voir dans la Figure 5.4 , trois fortes corrélations ont été révélées. Ces corrélations sont entre les valeurs max_{SNR} et min_{SNR} (valeur de corrélation = 1) et entre les valeurs max_{SNR} et min_{SNR} avec la saison (valeur de corrélation = 0,9). Des corrélations beaucoup plus faibles ont été trouvées entre les autres attributs.

Par la suite, un seuil arbitraire a été fixé à 0,7 comme valeur de corrélation dépendante. Cette valeur de seuil a été choisie de façon arbitraire sur la base que les valeurs de corrélation supérieure à 0,7 sont considérées comme forte (Moore, Notz et Fligner, 2013; Ratner, 2017). Ainsi, au-dessus de ce seuil, si deux attributs sont corrélés l'un à l'autre, un seul des deux attributs est conservé et l'autre est supprimé de l'ensemble des attributs. Les valeurs minimales et maximales (max_{SNR} et min_{SNR}) ont donc été supprimées, car en plus d'être corrélées entre elles, elles sont également corrélées respectivement avec la saison.

La deuxième étape de l'ingénierie des attributs est l'implémentation des méthodes de sélection des attributs. Ces méthodes utilisent trois algorithmes d'apprentissage machine couramment utilisés dans la littérature (Chandrashekar et Sahin, 2014; Degenhardt, Seifert et Szymczak, 2017; Stańczyk et Jain, 2015), à savoir l'algorithme Boruta, la méthode RF et la méthode d'élimination récursive des attributs (*Recursive Feature Elimination, RFE*). Les 3 algorithmes sont entraînés et testés à l'aide du package *Scikit-Learn* dans Python 3.

5.4.1.2 Sélection des attributs

Après la phase de nettoyage des attributs, l'étape finale de l'ingénierie des attributs est la sélection des attributs les plus pertinents pour le modèle. En effet, contrairement aux attributs redondants qui n'apportent aucune information supplémentaire au modèle, les attributs non pertinents peuvent ajouter du bruit et des erreurs dans le modèle. Cela peut causer ainsi une réduction de performance du modèle. Les méthodes de sélection d'attributs choisis (algorithme de Boruta, RF et RFE) permettent donc d'évaluer la pertinence de chaque attribut par rapport aux étiquettes définies.

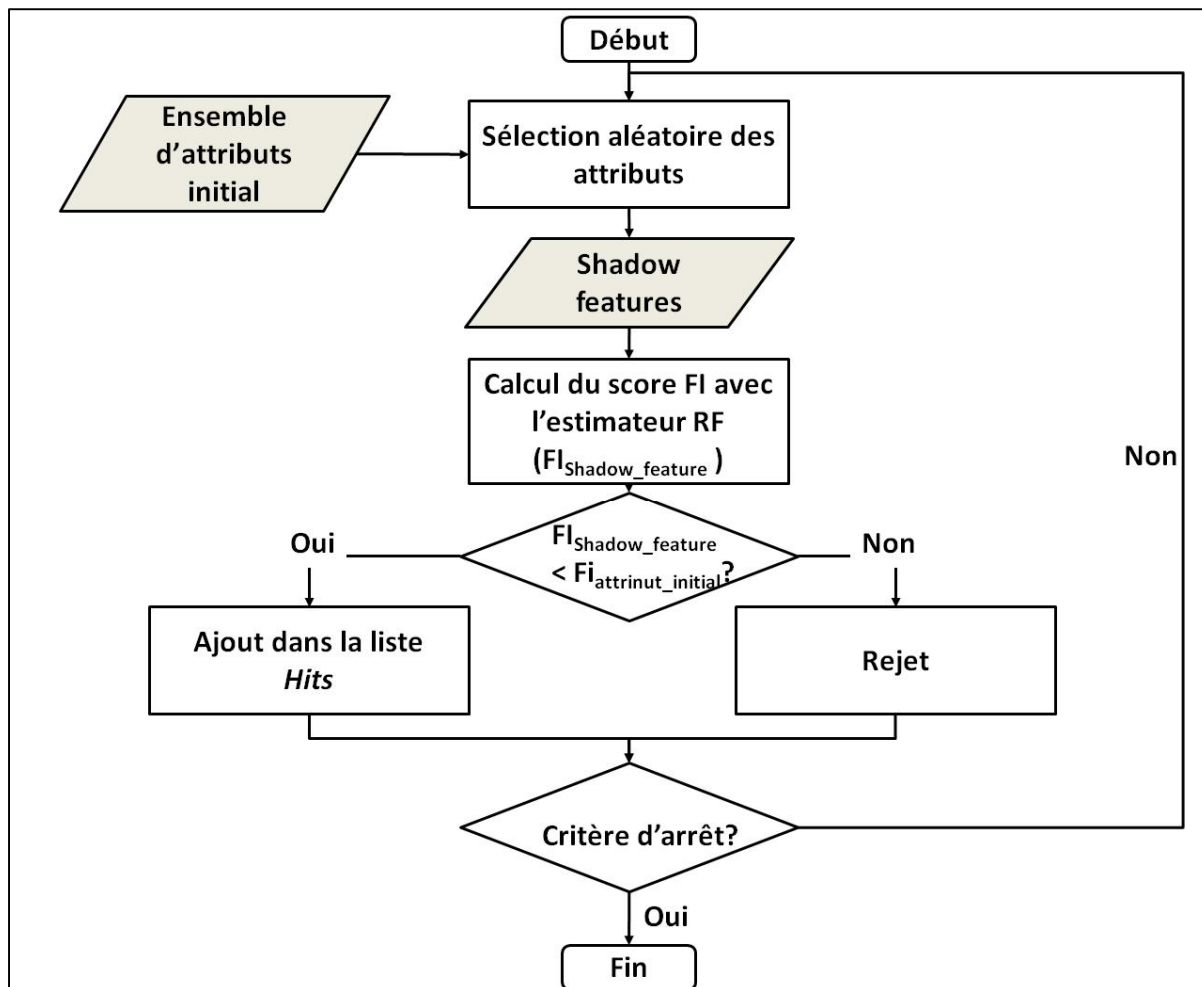


Figure 5.5 Processus de la méthode Boruta

- Algorithme Boruta : l'objectif est de sélectionner les attributs ayant le plus d'impact sur la sortie (étiquette) désirée. Ici, la sortie désirée est l'anomalie à détecter. Il a été implémenté en suivant les étapes décrites à la Figure 5.5. La première étape consiste à créer un nouvel ensemble d'attributs, appelés attributs d'ombre ou *Shadow feature*. Cet ensemble d'attributs est issu d'une combinaison tirée de l'ensemble d'attributs initial. La finalité de cette étape est de pouvoir comparer l'importance ou le score (*Feature Importance*, FI) des attributs du *Shadow feature* avec l'importance des attributs de l'ensemble d'attributs initial. La deuxième étape consiste à déterminer le score FI des attributs à partir d'un estimateur

entraîné sur chaque ensemble d'attributs générés. L'estimateur utilisé dans le cadre de notre étude est l'estimateur RF. À la fin de cette étape, le score FI de chaque attribut de la liste initial est comparée à la valeur maximale du score FI de tous les attributs du *Shadow feature*. Cela permet de créer une liste appelée *hits*. Cette liste contient tous les attributs ayant un score FI plus élevé que celui du meilleur attribut de l'ensemble *Shadow feature*. Finalement, la dernière étape est la répétition des deux étapes précédentes jusqu'à ce que tous les attributs aient été inscrits ou rejetés de la liste *hits* ou jusqu'à ce qu'une condition prédéfinie soit atteinte.

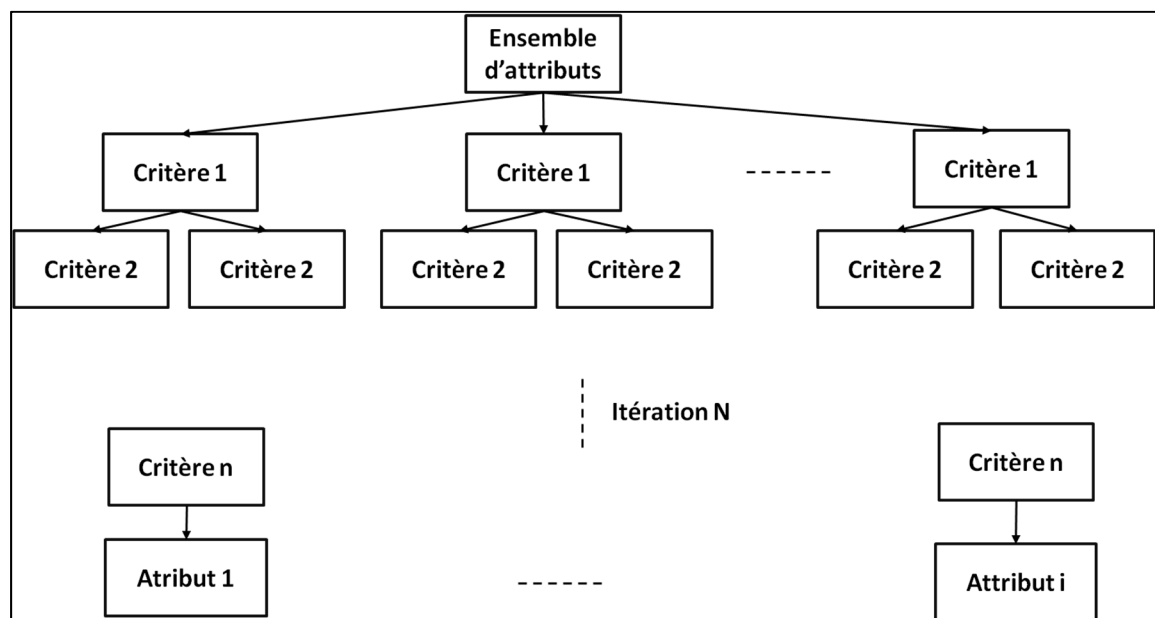


Figure 5.6 Processus de la méthode RF

- Méthode de forêt aléatoire (RF) : cette méthode est basée sur les arbres de décision. L'idée principale est de construire un arbre à partir de la base de données dans lequel chaque nœud de l'arbre est une condition sur l'attribut à insérer dans l'arbre. La Figure 5.6 présente un exemple d'arbre de décision. Ainsi, pour construire l'arbre de décision, la première étape est la détermination pour chaque nœud des mesures d'importance des attributs, appelées ici impuretés. Les attributs pertinents pour la classification auront une valeur d'impureté

élevée, alors que ceux qui ne sont pas associés aux sorties désirées auront une valeur d'impureté proche de zéro. Par la suite, cette mesure d'impureté est utilisée comme condition pour diviser la base de données en un groupe distinct formant les branches de l'arbre. La branche de l'arbre est donc formée de manière à ce que les attributs choisis soient ceux qui conduisent à une diminution de la mesure d'impureté. En d'autres termes, les attributs sélectionnés au sommet de l'arbre sont généralement plus importants.

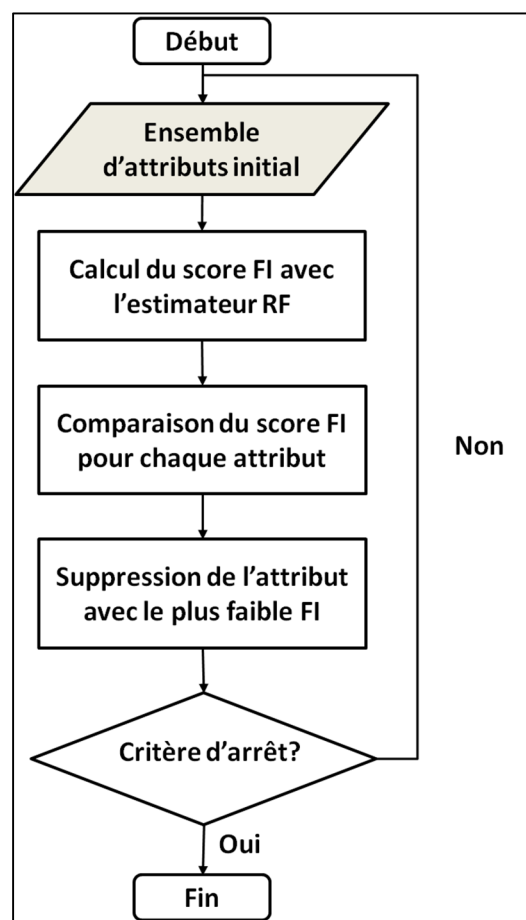


Figure 5.7 Processus du modèle RFE

- Méthode d'élimination récursive des attributs (RFE) : elle a pour objectif de trouver le plus petit ensemble d'attributs pour lequel le modèle fournit de bonnes performances. Tel que présenté à la Figure 5.7, cette méthode suit trois étapes. La première étape est la

détermination de l'importance de chaque attribut en utilisant un estimateur RF entraîné sur l'ensemble initial des attributs. La deuxième étape est l'étape d'élimination qui consiste à élaguer de l'ensemble d'attributs initial les attributs les moins importants. La dernière étape consiste à répéter les étapes précédentes de manière récursive sur l'ensemble élagué jusqu'à ce que le nombre souhaité d'attributs à sélectionner soit finalement atteint.

Tableau 5.1 Présentation des meilleurs attributs sélectionnés à partir des méthodes de sélection d'attributs

Boruta	RF	RFE
Saison	Saison	Résidu
Tendance	Tendance	
Résidu	Résidu	
Puissance optique au récepteur	Numéro de la semaine	
Jour		
Heure		
Numéro de la semaine (dans l'année)		

Le Tableau 5.1 présente les attributs trouvés pour chaque méthode de sélection des attributs. Ainsi, pour la méthode utilisant l'algorithme Boruta, les meilleurs attributs trouvés par cette méthode de sélection sont : saison, tendance, résidu, P_{RX} , jour, heure, $ind_{semaine}$. Pour la méthode utilisant le RF, les meilleurs attributs issus de la liste *hits* sont : saison, tendance, résidu et $ind_{semaine}$. Pour la méthode utilisant le RFE, le meilleur attribut en utilisant la méthode RFE est le résidu.

Après l'étape d'ingénierie des attributs, les attributs trouvés sont introduits dans le modèle de classification pour la détection anticipée des anomalies de SNR. On se retrouve donc avec quatre modèles de détection anticipée des anomalies SNR. Chacun de ces modèles possède un

ensemble différent d'attributs correspondant aux trois ensembles d'attributs trouvés par les modèles de sélection d'attributs et à l'ensemble de caractéristiques résultant de l'étude de corrélation. Ils utilisent également l'algorithme SVM pour leur implémentation.

5.4.2 Modèles de détection des anomalies

L'algorithme SVM a été utilisé pour implémenter les modèles de détection d'anomalies de SNR. En effet, régulièrement utilisé pour les problèmes de détection des anomalies, il surpasse plusieurs algorithmes d'apprentissage machine en offrant un bon compromis entre le temps de calcul et la précision des modèles (Braei et Wagner, 2020; Chandola, Banerjee et Kumar, 2009; Shahkarami et al., 2018; Singh, Thakur et Sharma, 2016).

Quatre modèles utilisant l'algorithme SVM ont été implémentés. Ces quatre modèles utilisent les attributs trouvés dans la section 5.4.1.2. Ils sont identifiés par SVM_all, SVM_Boruta, SVM_RF et SVM_RFE correspondant respectivement au modèle utilisant l'ensemble d'attributs après la phase de nettoyage, l'ensemble d'attributs Boruta, l'ensemble d'attributs RF et l'ensemble d'attributs RFE. Cette section inclut la présentation des modèles de détection d'anomalies et est divisée en deux parties. La première partie décrit le fonctionnement de l'algorithme SVM et son implémentation dans les modèles de détection d'anomalies. La deuxième partie est l'évaluation des performances des modèles.

5.4.2.1 Implémentation de l'algorithme SVM

L'algorithme SVM avait été présenté dans la section 2.3 et a pour but est de distinguer les classes de la base de données en utilisant un hyperplan. Cet hyperplan agit comme une frontière entre les classes. Dans notre cas, nous avons deux classes : la classe 0 (non anomalie) et la classe 1 (anomalie), telles que présentées dans la section 5.3.2. Pour rappel, le fonctionnement du SVM est présenté à la Figure 5.8.

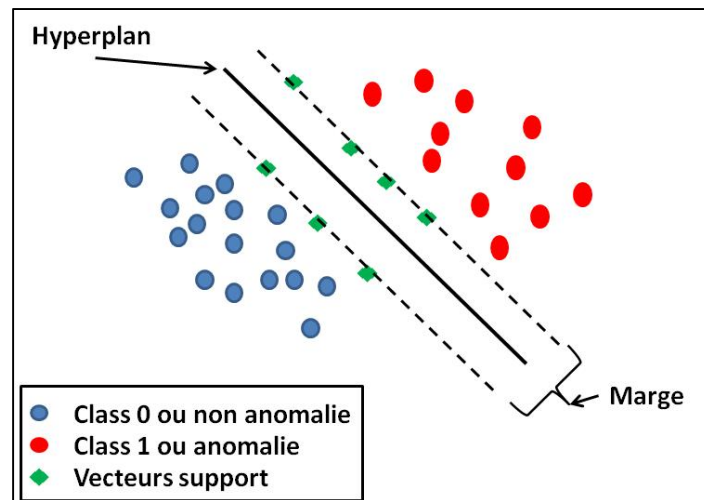


Figure 5.8 Fonctionnement général de l'algorithme SVM

L'hyperplan peut être une frontière linéaire ou non linéaire. Il est construit de manière à maximiser la distance entre les deux classes. La distance entre les classes s'appelle la marge. L'emplacement de la marge est déterminé à l'aide d'un sous-ensemble d'observations pris dans l'ensemble d'apprentissage. Ce sous-ensemble d'observations est appelé vecteurs de support.

Notons que l'algorithme SVM est défini par plusieurs hyper-paramètres, à savoir :

- la régulation, l'équilibrage ou le paramètre C : ce paramètre permet d'éviter ou de réduire les erreurs de classification dans les ensembles d'apprentissage. Plus la valeur de C est élevée, plus la marge est petite avec une classification plus stricte. À l'inverse, plus la valeur de C est petite, plus la marge est grande avec une classification plus large ;
- le paramètre gamma ou G : il définit l'influence des points d'entraînement (ou point d'apprentissage) sur le calcul de l'hyperplan. En d'autres termes, un paramètre G bas signifie que les points définissant les vecteurs de support pourraient être éloignés de l'hyperplan. De même, un paramètre G élevé signifie que seuls les points les proches de l'hyperplan sont considérés pour définir les vecteurs de support ;
- le noyau : il représente la fonction qui définit l'ensemble de représentation des attributs. Il est utilisé pour résoudre l'équation de l'hyperplan

Ainsi, afin d'implémenter les quatre modèles de détection anticipée des anomalies de SNR, la base de données est divisée selon le ratio 80/20. Il correspond respectivement aux jeux de données d'entraînement et de test. Le jeu de données d'entraînement est utilisé pour optimiser les hyper-paramètres à utiliser, à savoir, les paramètres C et G, dans les modèles et aussi pour entraîner les modèles. Pour ce faire, la méthode *GridSearchCV* de la bibliothèque *Scikit-Learn* a été utilisée. Cette méthode consiste à effectuer une recherche exhaustive sur une combinaison de valeurs d'hyper-paramètres afin de trouver la meilleure combinaison à utiliser. Le noyau, quant à lui, a été fixé. En effet, parce que la relation entre les anomalies et les attributs est non linéaire, celui-ci a été fixé à la fonction de base radiale (*Radial Basis Function, RBF*) (Hsu, Chang et Lin, 2008; Shmilovici, 2005). Le Tableau 5.2 présente les hyper-paramètres optimaux trouvés avec la méthode *GridSearchCV*.

Tableau 5.2 Ensemble des hyper-paramètres optimaux

	C	G	Noyau
SVM_all 8 attributs	1000	0.1	RBF
SVM_Boruta 7 attributs	1000	0.1	RBF
SVM_RF 4 attributs	1000	1	RBF
SVM_RFE 1 attribut	1000	1	RBF

Les valeurs utilisées sont (0,1 ; 1 ; 10 ; 100 ; 1000) pour le paramètre C et (1 ; 0,1 ; 0,001 ; 0,001 ; 0,0001) pour le paramètre G. La combinaison finale (C ; G) est (1000 ; 0,1) pour les modèles SVM_all et SVM_Boruta. Pour les modèles SVM_RF et SVM_RFE, la combinaison obtenue est (1000 ; 1).

5.4.2.2 Évaluation des performances

L'évaluation des modèles a été effectuée en utilisant le jeu de test correspondant à 20% de la base de données. Celui-ci inclut 2 108 instances de classe 0 (non-anomalie) et 410 instances de classe 1 (anomalie).

La Figure 5.9 montre les matrices de confusion pour chaque modèle.

		(a)		(b)		(c)		(d)	
		Oui	Non	Oui	Non	Oui	Non	Oui	Non
Classe prédite	Oui	407	9	408	8	406	6	410	4
	Non	3	2099	2	2100	4	2102	0	2104
		Classe réelle							

Figure 5.9 Matrices de confusion pour les modèles : (a) SVM_All ; (b) SVM_Boruta ; (c) SVM_RF ; (d) SVM_RFE

Tirée de Allogba, Yaméogo et C. (2021)

Celles-ci donnent un aperçu de la précision des modèles à partir des taux de vrais positifs et faux positifs situés dans la diagonale des matrices. Le taux de vrais positifs correspond aux bonnes prédictions faites par les modèles sur l'ensemble des 2 518 instances de l'ensemble de données de test. Plus celui-ci est élevé, meilleur est le modèle. Quant au taux de faux positifs, il indique le cas où le modèle prédit une non-anomalie (classe 0) alors que la donnée réelle est une anomalie (classe 1). Contrairement au taux de vrais positifs, plus le taux de faux positif est faible, meilleur est le modèle. Avec son taux de faux positifs le plus bas (0) et son taux de vrais positifs le plus élevé (410), le modèle SVM_RFE est le meilleur modèle, suivi par le modèle SVM_Boruta.

Par ailleurs, les matrices de confusion sont également utilisées pour déterminer les métriques de performance des modèles. Celles-ci sont le score F1 pour évaluer la précision des modèles, le rappel pour déterminer la proportion d'anomalies correctement prédites parmi toutes les anomalies observées et le temps de calcul pour comparer les temps d'exécution des modèles. Elles sont présentées dans le Tableau 5.3. Notons que le score F1 et le rappel sont étudiés en cas de déséquilibre dans les données, c'est-à-dire lorsque la proportion d'une classe est largement supérieure à une autre. Dans la présente étude, constituant 83,66 % de la base de données de test, le nombre d'instances dans la classe 0 (non-anomalie) est supérieur au nombre d'instances de la classe 1 (anomalie).

Tableau 5.3 Présentation des performances des modèles de détection des anomalies

	Score F1 (%)	Rappel (%)	Temps de calcul (ms)¹
SVM_all 8 attributs	98,53	99,26	46,87
SVM-Boruta 7 attributs	98,78	99,51	46,87
SVM-RF 4 attributs	98,78	99,02	15,62
SVM-RFE 1 attribut	99,51	100	5

¹ Les modèles ont été exécutés sur un système Intel® Core™

i5-7200U 2,5 GHz CPU, avec 8 GB de mémoire RAM

Notons que les modèles ont été exécutés une seule fois dans le jeu de test afin d'évaluer les performances. Une analyse statistique sur les métriques trouvées reste à être effectuée afin d'évaluer la robustesse des modèles.

Le score F1 montre la qualité du modèle. Plus la valeur du score F1 est élevée, meilleur est le modèle. En observant le Tableau 5.3, les meilleures performances, en termes de score F1, ont

été obtenues avec le modèle de détection des anomalies SVM_RFE. Ce modèle correspond au modèle ayant le plus petit nombre d'attributs. En d'autres termes, il représente le meilleur modèle pour prédire correctement les anomalies de SNR. En effet, le score F1 des modèles diminue de 99,51% pour le modèle SVM_RFE (1 attribut) à 98,78%, 97,78% et 98,53% respectivement pour les modèles SVM_Boruta (7 attributs), SVM_RF (4 attributs) et SVM_all (tous les 8 attributs). Notons également que plus le nombre d'attributs est petit, meilleure est la performance. Comme présentées par Chandrashekar et Sahin (2014), les méthodes d'ingénierie des attributs permettent d'améliorer les performances des modèles de prédiction en réduisant les attributs non pertinents pouvant générer des erreurs dans la prédiction. Cela explique donc le fait que les modèles avec un nombre réduit d'attributs performant mieux que ceux avec un nombre plus grand d'attributs.

Par ailleurs, un bon modèle de détection des anomalies signifie de faibles taux de faux positifs observés dans la matrice de confusion (Figure 5.9). Cela équivaut à un rappel élevé. En effet, le taux de faux positif représente les cas où le modèle prédit une instance comme non-anomalie (classe 0) alors que la vraie classe de l'instance est étiquetée comme une anomalie (classe 1). Ainsi, en analysant le rappel présenté dans le Tableau 5.3, le meilleur modèle est le modèle SVM_RFE avec un rappel à 100%. Cela signifie donc que le modèle SVM_RFE détecte toutes les anomalies présentes dans la base de test. Concernant les autres modèles, le modèle SVM_Boruta donne un meilleur rappel avec 99,51%, suivi du modèle SVM_all avec un rappel de 99,26% et du modèle SVM_RF avec un rappel de 99,02%.

Enfin, concernant le temps de calcul, le nombre d'attributs affecte la rapidité de prédiction du modèle. Moins les modèles utilisent d'attributs, plus rapide est le modèle. Le temps de calcul explique également la complexité du modèle. En d'autres termes, plus le nombre d'attributs est faible, moins complexe est le modèle. Ainsi, bien qu'ayant des métriques de performance (score F1 et rappel) sensiblement les mêmes, le modèle SVM_RFE se distingue avec le plus petit temps de calcul, ce qui le positionne comme étant le modèle le moins complexe et le plus intéressant à utiliser.

En outre, en comparant les attributs trouvés dans chacun des modèles, l'attribut commun à tous les modèles est le résidu. Ceci valide donc l'hypothèse émise dans Yaméogo et al. (2020) dans lequel le résidu se présente comme le meilleur attribut pour caractériser les anomalies de SNR.

Le Tableau 5.4 présente sous forme de synthèse une analyse comparative de performance du modèle de détection des anomalies proposé et des principales méthodes développées à ce jour pour la détection des anomalies.

Tableau 5.4 Synthèse des résultats

Références	Méthode	Base de données	Taux de succès
Vela et al. (2017a; 2017b; 2017c) (Méthode 1)	<u>Définition des anomalies :</u> Méthode par seuil fixe <u>Attributs utilisés :</u> Valeurs statistiques <u>Détection des anomalies</u> Modèle à état fini (<i>Finite State Machine</i> , FSM)	<u>Types de données :</u> Données synthétiques BER <u>Longueur de la connexion :</u> 320 km <u>Période d'observation :</u> 60 jours	90%
Shahkrami et al. (2018) (Méthode 2)	<u>Définition des anomalies :</u> Méthode par seuil fixe <u>Attributs utilisés :</u> Valeurs statistiques <u>Détection des anomalies</u>	<u>Types de données :</u> Données synthétiques BER <u>Longueur de la connexion :</u> 380 km	98.4% 99% 99.1%

Références	Méthode	Base de données	Taux de succès
	NN SVM RF	<u>Période d'observation</u> : 24 heures	
Chen et al. (2019; 2018) (Méthode 3)	<u>Définition des anomalies</u> : <i>Clustering</i> <u>Attributs utilisés</u> : Densité du <i>cluster</i> <u>Détection des anomalies</u> : NN	<u>Types de données</u> : Données synthétiques Puissance <u>Longueur du réseau</u> : 220 km – 6 connexions <u>Période d'observation</u> : Non défini	94.8%
Allogba et al.	<u>Définition des anomalies</u> Méthode IQR <u>Attributs utilisés</u> : Valeurs statistiques Puissance Données de la décomposition en série temporelle <i>Feature engineering</i> <u>Détection des anomalies</u> : SVM	<u>Types de données</u> : Données réelles SNR <u>Longueur de la connexion</u> : 150 à 1250 km 8 connexions <u>Période d'observation</u> : 24 heures	100%

Comme présenté dans le Tableau 5.4, le modèle proposé pour la détection des anomalies a donné de meilleur performant en utilisant le modèle SVM_RFE. En effet, avec 100% de précision (100% de succès) comparativement à 90%, 99,1% et 98,4% pour les méthodes 1, 2 et 3 publiées dans la littérature.

De plus, en utilisant des données réelles issues d'un réseau réelle de production avec des connexions allant jusqu'à 1250 km, contrairement aux autres méthodes utilisant des données synthétiques, il a été possible d'effectuer une ingénierie des attributs. Celle-ci a permis d'extraire des 8 connexions le résidu, comme meilleure entrée à utiliser dans l'algorithme SVM, comparativement aux valeurs statistiques utilisées dans les méthodes 1 et 2 de la littérature.

5.5 Conclusion

Dans ces travaux, nous avons proposé un nouveau modèle pour la détection anticipée des anomalies de SNR. Plus précisément, le modèle de détection permet d'examiner les instances de SNR une à une et d'en extraire le meilleur attribut. Par la suite, l'algorithme SVM est utilisé pour classer cette nouvelle instance en tant qu'anomalie ou pas. Ce modèle est ainsi construit à partir de deux modules distincts décrits ci-après.

Le premier module offre une nouvelle définition des anomalies basée sur la technique statistique IQR. Cette technique a permis d'identifier dans les séries temporelles de SNR de 8 connexions optiques d'un réseau de production des données qui pourraient être considérées comme des anomalies de SNR. Il se distingue également de la méthode traditionnelle, coûteuse pour les opérateurs, qui consiste à fixer un seuil au-delà duquel l'anomalie est détectée. De plus, ce module a permis d'identifier les attributs à utiliser pour détecter les anomalies. Ceux-ci ont été regroupés en trois catégories, à savoir les métriques de performance collectées (puissance optique au récepteur, et les limites des valeurs aberrantes des données de SNR), les caractéristiques temporelles (jour, heure, la semaine dans l'année, la période de la journée) et

les caractéristiques statistiques issues des décompositions temporelles des données de SNR (saison, tendance, résidu). Cela diffère des attributs proposés dans la littérature, principalement composés des paramètres statistiques (valeurs minimales, valeurs maximales, etc.).

Le second module permet de détecter les anomalies dans les données de SNR. Il est effectué premièrement par une ingénierie des attributs qui a pour but de réduire la liste des attributs à ceux qui sont les plus pertinents. Celle-ci est réalisée tout d'abord par une analyse de corrélation pour supprimer les attributs redondants. Trois algorithmes d'apprentissage machine pour la sélection des attributs ont ensuite été implémentés : le RF, la méthode Boruta et le RFE. L'ingénierie des attributs a permis de réduire le nombre initial d'attributs (10 attributs) à 4, 7 et 1 attributs respectivement pour le RF, la méthode Boruta et le RFE. Ensuite, à partir des attributs sélectionnés, quatre modèles de détection des anomalies utilisant l'algorithme SVM ont été proposés. Cet algorithme permet à partir des différentes combinaisons des attributs de mesurer leur impact sur les anomalies et de détecter la présence d'anomalies ou pas à dans les données. Les meilleurs résultats pour la détection des anomalies ont été obtenus en utilisant le modèle SVM-RFE (un seul attribut) avec un score F1 de 99,51%. L'analyse de l'impact des attributs a également montré que plus le nombre d'attributs est réduit, meilleure est la détection. En effet, la précision des modèles en regardant le score F1 diminue en fonction du nombre d'attributs. De plus, en analysant les performances des différents modèles, notre approche a également permis de montrer que le résidu est le meilleur attribut, en comparaison des autres utilisés, pour caractériser les anomalies de SNR.

Le modèle SVM-RFE, utilisant le résidu comme attribut, avec un temps de calcul très bas, se présente comme une solution qui pourrait être utilisée en temps réel pour détecter les anomalies de SNR. Il pourrait être implémenté dans des outils de maintenance proactive et ainsi offrir une solution de maintenance proactive en cas de panne dans le réseau due à une dégradation de performance du réseau.

CONCLUSION

L'arrivée de la cognition dans le domaine de la télécommunication a permis d'utiliser les métriques de performance du réseau afin d'implémenter des outils proactifs. En effet, la demande grandissante des applications a augmenté le besoin d'avoir des outils de gestion et de contrôle capable de répondre à l'évolutivité et au dynamisme du réseau. Ainsi, les travaux présentés dans cette thèse ont porté sur le développement de nouvelles méthodes, basées sur les algorithmes d'apprentissage machine, pour l'estimation et la prédiction de la qualité de transmission des connexions optiques déployées dans un réseau, ainsi que pour la détection des anomalies de qualité de transmission. Notons que les récents travaux effectués dans ce sens ont utilisé pour la plupart des données synthétiques pour l'implémentation de leurs méthodes, ne reflétant généralement pas les événements réels observés sur le terrain dans des réseaux de production. Les modèles présentés dans cette thèse ont été construits à partir de méthodes existantes dans la littérature mais entraînés et testés pour la première fois en utilisant des données de performance collectées sur des connexions optiques déployées dans des réseaux de production.

Des outils permettant l'estimation de la performance d'une connexion optique ont été proposés. Cela a d'abord été effectué par un outil de classification des données de BER d'une connexion optique en utilisant les propriétés statistiques des données historiques recueillies (valeurs maximales et moyennes) ainsi que les facteurs externes comme la température externe, les activités humaines définies par la période de la journée. Il a été ainsi possible de classifier les données de BER avec un taux de précision de 97,8%.

Une ingénierie des attributs a également été implémentée dans deux modèles d'estimateur de la qualité d'une connexion optique à établir. Ces modèles utilisent comme attributs les caractéristiques des connexions optiques, à savoir la longueur de la liaison, le gain et le nombre d'amplificateurs optiques, de même que le format de modulation, la puissance et le débit du canal. Cette ingénierie des attributs avait pour but d'évaluer la possibilité d'estimer la qualité

d'une nouvelle connexion optique à établir dans un contexte réel dans lequel toutes les données pouvaient ne pas être disponibles. L'ingénierie des attributs a permis de valider la performance des deux modèles d'estimation de la qualité d'une connexion optique dans le cas où le nombre d'attributs était réduit, avec un taux de précision allant de 93.30% (cas avec 4 attributs) à 88.52% (cas avec 3 attributs).

Ensuite, des algorithmes d'apprentissage machine ont été explorés pour la prédiction du SNR de connexions optiques sur un horizon jusqu'à 96 heures à partir de données historiques recueillies dans des réseaux de production. Les algorithmes d'apprentissage machine utilisés dans cette étude sont basés sur les réseaux de neurones. Cela offre ainsi une robustesse des modèles face aux effets de non stationnarité ou non saisonnalité des séries temporelles. Par ailleurs, ces modèles ont été à la fois construits et testés avec des données de terrain pour la première fois.

En outre, la prédiction du SNR a été effectuée en trois parties. La première partie a été l'implémentation de modèles univariés LSTM et GRU pour la prédiction du SNR en suivant trois scénarios :

- dans le premier scénario, le modèle LSTM-U-1A a été implémenté à partir des données de SNR d'une connexion optique A collectées pendant une période de 13 mois. Il a été possible de prédire les données de SNR pour un horizon de prévision allant jusqu'à 24 h par rapport à une méthode naïve, avec une différence de RMSE de 0,05 dB ;
- dans le deuxième scénario, le modèle GRU-U-1B a été implémenté à partir des données de SNR collectées sur une connexion optique B pendant une période de 18 mois. Il a été possible cette fois-ci de prédire les données de SNR pour des horizons allant jusqu'à 48 heures. Toutefois, les performances observées ont montré que le modèle GRU-U-1B performe mieux qu'un modèle naïf pour des horizons de prédiction inférieurs à 12 heures ;
- dans le troisième scénario, une étude comparative a été effectuée entre les modèles GRU et LSTM (appelés respectivement GRU-U-2A et LSTM-U-2A) en utilisant les données de

SNR de la connexion optique A. De nouveaux hyper-paramètres ont été pris en compte dans l'implémentation de ces modèles. Ainsi, il a été possible de prédire les données de SNR à partir de ces modèles pour des horizons allant jusqu'à 96 heures. De plus, les modèles LSTM-U-2A et GRU-U-2A ont donné de meilleures performances qu'un modèle naïf et un modèle encodeur-décodeur LSTM.

La deuxième partie a regroupé l'implémentation de deux modèles multivariés LSTM et GRU (appelés respectivement LSTM-M-1 et GRU-M-1) pour la prédiction du SNR de la connexion optique A. Les attributs utilisés pour ces modèles sont les métriques de performance collectées sur la liaison (pré-FEC BER, la puissance reçue du canal, le DGD), la température et la période de la journée. Il a été ainsi possible de prédire les nouvelles données de SNR pour des horizons de prédiction allant jusqu'à 96 heures. De plus, en comparant les modèles multivariés avec les modèles univariés GRU-U-2A et LSTM-U-2A ainsi qu'une méthode naïve, les modèles multivariés ont montré une légère amélioration des performances, avec des valeurs RMSE inférieures des valeurs de R^2 plus élevées. Par ailleurs, une étude d'ingénierie d'attributs a été effectuée dans le but de réduire les attributs à utiliser dans le modèle. Cette étude a révélé qu'avec la réduction des attributs, il était toujours possible de prédire les valeurs de SNR pour des horizons allant jusqu'à 96 heures. De plus, elle a également montré que parmi les attributs utilisés, la puissance du canal était le meilleur attribut additionnel (parmi les attributs disponibles) à utiliser avec les valeurs passées de SNR.

La troisième partie a été l'exploration du concept d'apprentissage par transfert. Cette partie a permis de montrer qu'il a été possible de prédire les données de SNR de nouvelles connexions optiques en utilisant le modèle multivarié avec le nombre d'attributs réduits LSTM-M-2. Notons que pour les deux connexions, l'une des connexions appartenait à la même liaison que le domaine source tandis que l'autre appartenait à une liaison différente du domaine source.

Finalement, un modèle de détection des anomalies de SNR utilisant l'algorithme SVM a été implémenté. Pour ce faire, la définition et l'extraction des anomalies de SNR ont été

premièrement effectuées en utilisant une méthode basée sur la technique statistique IQR. Notons également que les attributs caractérisant les anomalies et utilisés pour détecter les anomalies de SNR sont issues des métriques de performance collectées (puissance optique au récepteur P_{RX}), des caractéristiques temporelles (jour, heure, semaine de l'année ($ind_{semaine}$), période de la journée($période_{jour}$)), des caractéristiques statistiques issues des décompositions temporelles des données de SNR (saison, tendance, résidu) et des limites supérieures et inférieures des valeurs aberrantes des données de SNR. Dans un deuxième temps, une ingénierie des attributs a été effectuée utilisant trois techniques d'apprentissage machine : le Boruta, le RF et le RFE. Celle-ci a permis de réduire le nombre initial d'attributs à 7, 4 et 1 attribut, respectivement. À partir des attributs sélectionnés, il a donc été possible de prédire les anomalies de SNR en utilisant le modèle SVM-RFE (contenant un seul attribut) avec un score F1 de 99,51%.

Les différents modèles proposés représentent ainsi des solutions prometteuses pouvant être implémentées dans les outils de gestion et de contrôle des réseaux optiques. Ils permettraient ainsi aux opérateurs de déployer des actions proactives face à une panne prédite par les modèles.

RECOMMANDATIONS

Afin de continuer l'exploration de cette thèse, cinq pistes d'amélioration sont proposées :

- estimation de la QoT : dans le cadre de l'estimation de la qualité de transmission d'une connexion optique avant son établissement, l'une des pistes d'amélioration est l'exploration d'autres techniques d'apprentissage machine, principalement les techniques basées sur des ensembles de classificateurs (*classifier pools*). Ces techniques permettraient d'effectuer une combinaison des modèles de classification de base afin d'améliorer les performances de classification ;
- amélioration de la méthode d'apprentissage par transfert pour la prédiction des données de SNR : la piste de la méthode d'apprentissage par transfert se présente comme une solution prometteuse pour les opérateurs dans le sens qu'elle permettrait d'utiliser un seul modèle pour la prédiction de données de SNR de différentes connexions. Dans ce contexte, une première étude pourrait être faite sur l'optimisation des hyper-paramètres utilisés dans notre modèle d'apprentissage par transfert. Rappelons que celui-ci n'avait pas été optimisé. Une deuxième étude pourrait être d'utiliser un plus grand nombre de données dans le domaine source. Pour ce faire, la première étape serait de collecter et de classer les connexions optiques des domaines sources en fonction de plusieurs critères, tels que la distribution des données de SNR, la liaison commune des connexions, etc. Par la suite, une seconde étape serait d'implémenter différents modèles de prédiction, issus de chaque groupe classifié, et de les tester dans les différents domaines sources. Cela aurait pour but d'analyser les performances des modèles pour chaque donnée de SNR ;
- amélioration des modèles multivariés : les modèles multivariés se sont présentés comme étant des modèles prometteurs pour les situations temps réel de par leur performance et leur temps de prédiction élevée. Une piste d'amélioration de ces modèles serait l'optimisation des hyper-paramètres, en utilisant une plage plus grande de valeurs des hyper-paramètres ainsi que les techniques d'optimisation automatique des paramètres, telles que présentées par Wu et al. (2019). Notons que l'optimisation automatique des hyper-paramètres pourrait également se faire pour les modèles univariés.

- détection des anomalies : dans le cadre de la prédiction des anomalies, l'une des pistes d'amélioration est l'exploration d'autres techniques d'apprentissage machine. Celle-ci a pour but à la fois d'évaluer et de comparer notre modèle développé face à d'autres modèles et également de pouvoir évaluer la fiabilité des modèles de détection d'anomalies en temps réel ;
- évaluation des performances : lors de l'évaluation des performances des différents modèles, ces derniers ont été exécutés une seule fois. Afin d'évaluer la robustesse des modèles, l'une des pistes d'amélioration serait d'effectuer des analyses statistiques sur les métriques étudiées. Ces analyses statistiques permettraient de déterminer les intervalles de confiance pour chacune des métriques et d'effectuer une comparaison plus solide entre les différents modèles implémentés.

LISTE DE RÉFÉRENCES BIBLIOGRAPHIQUES

- Aladin, S., A. V. S. Tran, S. Allogba et C. Tremblay. 2020a. « Quality of Transmission Estimation and Short Term Performance Forecast of Lightpaths ». *Journal of Lightwave Technology*, vol. 38, n° 10, p. 2807-2814.
- Aladin, Sandra, Stéphanie Allogba, Anh Vu Stephan Tran et Christine Tremblay. 2020b. « Recurrent Neural Networks for Short-Term Forecast of Lightpath Performance ». In *Optical Fiber Communication Conference (OFC) 2020*. (San Diego, California, 2020/03/08), p. W2A.24. Coll. « OSA Technical Digest »: Optical Society of America. < <http://www.osapublishing.org/abstract.cfm?URI=OFC-2020-W2A.24> >.
- Aladin, Sandra, et Christine Tremblay. 2018. « Cognitive Tool for Estimating the QoT of New Lightpaths ». In *Optical Fiber Communication Conference*. (San Diego, California, 2018/03/11), p. M3A.3. Coll. « OSA Technical Digest (online) »: Optical Society of America. < <http://www.osapublishing.org/abstract.cfm?URI=OFC-2018-M3A.3> >.
- Allogba, S., S. Aladin et C. Tremblay. 2021. « Multivariate Machine Learning Models for Lightpath Performance Forecasting ». *Journal of Lightwave Technology*
- Allogba, S., et C. Tremblay. 2018. « K-Nearest Neighbors Classifier for Field Bit Error Rate Data ». In *2018 Asia Communications and Photonics Conference (ACP)*. (26-29 Oct. 2018), p. 1-3.
- Allogba, S., et C. Tremblay. June 5-7, 2018. « Statistical analysis of performance monitoring data collected in an optical link ». In *20th Photonics North Conf.* (Montreal, Canada).
- Allogba, S., B. L. M. Yaméogo et Tremblay C. 2021. « Extraction and early detection of anomalies using Machine Learning models ». *Journal of Lightwave Technology*.
- Allogba, Stéphanie, et Christine Tremblay. June 5-7, 2017. « Prediction of BER Trends in an Optical Link Using Machine Learning ». In *19th Photonics North Conf.* (Ottawa, Canada, June 6-8, 2017).
- Amirabadi, M.A. 2019. « A Survey on Machine Learning for Optical Communication [Machine Learning View] ». *arXiv: Signal Processing*.
- Arslan, Hüseyin. 2007. *Cognitive Radio, Software Defined Radio, and Adaptive Wireless Systems*. Springer.

- Azodolmolky, S., J. Perelló, M. Angelou, F. Agraz, L. Velasco, S. Spadaro, Y. Pointurier, A. Francescon, C. V. Saradhi, P. Kokkinos, E. Varvarigos, S. A. Zahr, M. Gagnaire, M. Gunkel, D. Klonidis et I. Tomkos. 2011. « Experimental Demonstration of an Impairment Aware Network Planning and Operation Tool for Transparent/Translucent Optical Networks ». *Journal of Lightwave Technology*, vol. 29, n° 4, p. 439-448.
- Azodolmolky, Siamak, Mirosław Klinkowski, Eva Marin, Davide Careglio, Josep Solé Pareta et Ioannis Tomkos. 2009. « A survey on physical layer impairments aware routing and wavelength assignment algorithms in optical networks ». *Computer Networks*, vol. 53, n° 7, p. 926-944.
- Azzimonti, D., C. Rottondi et M. Tornatore. 2020. « Reducing probes for quality of transmission estimation in optical networks with active learning ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 12, n° 1, p. A38-A48.
- Balanici, M., et S. Pachnicke. 2019. « Machine Learning-Based Traffic Prediction for Optical Switching Resource Allocation in Hybrid Intra-Data Center Networks ». In *2019 Optical Fiber Communications Conference and Exhibition (OFC)*. (3-7 March 2019), p. 1-3.
- Bishop, Christopher M. 2006. *Pattern recognition and machine learning* (2006). New York: Springer, xx, 738 p. p.
- Borkowski, R., R. Duran, C. Kachris, D. Siracusa, A. Caballero, N. Fernandez, D. Klonidis, A. Francescon, T. Jimenez, J. Aguado, I. Miguel, E. Salvadori, I. Tomkos, R. Lorenzo et I. Monroy. 2015. « Cognitive optical network testbed: EU project CHRON ». *Optical Communications and Networking, IEEE/OSA Journal of*, vol. 7, n° 2, p. A344-A355.
- Braei, Mohammad, et Sebastian Wagner. 2020. « Anomaly Detection in Univariate Time-series: A Survey on the State-of-the-Art ». *ArXiv*, vol. abs/2004.00433.
- Brockwell, Peter J., et Richard A. Davis. 2002. *Introduction to time series and forecasting*, Second edition. New York: Springer. In WorldCat.org.
< <http://site.ebrary.com/id/10129778> >.
- Canada, Environnement et Changement climatique. 2011. « Données historiques - Climat - Environnement et Changement climatique Canada ». < https://climat.meteo.gc.ca/historical_data/search_historic_data_f.html >.
- Cassidy, Andrew Michael. 2012. « Performance dynamique de réseaux optiques sans filtre : ressources réseau et qualité de transmission ». Mémoire. École de Technologie Supérieure.

- Chan, Calvin C. K. 2010. *Optical Performance Monitoring - Advanced Techniques for Next-Generation Photonic Networks*. Elsevier.
- Chandola, Varun, Arindam Banerjee et Vipin Kumar. 2009. « Anomaly detection: A survey ». *ACM Comput. Surv.*, vol. 41, n° 3, p. Article 15.
- Chandrashekar, Girish, et Ferat Sahin. 2014. « A survey on feature selection methods ». *Computers & Electrical Engineering*, vol. 40, n° 1, p. 16-28.
- Chen, X., B. Li, R. Proietti, Z. Zhu et S. J. B. Yoo. 2019. « Self-Taught Anomaly Detection With Hybrid Unsupervised/Supervised Machine Learning in Optical Networks ». *Journal of Lightwave Technology*, vol. 37, n° 7, p. 1742-1749.
- Chen, X., B. Li, M. Shamsabardeh, R. Proietti, Z. Zhu et S. J. B. Yoo. 2018. « On Real-Time and Self-Taught Anomaly Detection in Optical Networks Using Hybrid Unsupervised/Supervised Learning ». In *2018 European Conference on Optical Communication (ECOC)*. (23-27 Sept. 2018), p. 1-3.
- Choudhury, G., D. Lynch, G. Thakur et S. Tse. 2018. « Two use cases of machine learning for SDN-enabled ip/optical networks: traffic matrix prediction and optical path performance prediction [Invited] ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 10, n° 10, p. D52-D62.
- Cristina Rottondi, Luca Barletta, Alessandro Giusti, Massimo Tornatore. 2018. « Machine-Learning Method for Quality of Transmission Prediction of Unestablished Lightpaths ». *Journal of Optical Communications and Networking*, vol. 10, n° 2.
- Dancho, Matt, et Davis Vaughan. 2018. « Package ‘anomalize’ ». < <https://github.com/business-science/anomalize> >.
- Degenhardt, Frauke, Stephan Seifert et Silke Szymczak. 2017. « Evaluation of variable selection methods for random forests and omics data sets ». *Briefings in Bioinformatics*, vol. 20, n° 2, p. 492-503.
- Fernández, N., R. J. Durán, I. de Miguel, N. Merayo, P. Fernández, J. C. Aguado, R. M. Lorenzo, E. J. Abril, E. Palkopoulou et I. Tomkos. 2013. « Virtual Topology Design and reconfiguration using cognition: Performance evaluation in case of failure ». In *2013 5th International Congress on Ultra Modern Telecommunications and Control Systems and Workshops (ICUMT)*. (10-13 Sept. 2013), p. 132-139.
- Francesco Musumeci, Cristina Rottondi, Avishek Nag, Irene Macaluso, Darko Zibar, Marco Ruffini et Massimo Tornatore. 2018. « A Survey on Application of Machine Learning Techniques in Optical Networks ».

- Gao, Ruoxuan, Lei Liu, Xiaomin Liu, Huazhi Lun, Lilin Yi, Weisheng Hu et Qunbi Zhuge. 2020. « An overview of ML-based applications for next generation optical networks ». *Science China Information Sciences*, vol. 63, n° 6, p. 160302.
- Gu, Rentao, Zeyuan Yang et Yuefeng Ji. 2020. « Machine learning for intelligent optical networks: A comprehensive survey ». *Journal of Network and Computer Applications*, vol. 157, p. 102576.
- Hsu, Chih-Wei, Chih-Chung Chang et C. Lin. 2008. « A Practical Guide to Support Vector Classification ». In.
- Hyndman, Rob J, et George Athanasopoulos. 2018. *Forecasting: principles and practice*, 2nd edition. Melbourne, Australia: OTexts.
- Jimenez, T., J. C. Aguado, I. de Miguel, R. J. Duran, N. Fernandez, M. Angelou, D. Sanchez, N. Merayo, P. Fernandez, N. Atallah, R. M. Lorenzo, I. Tomkos et E. J. Abril. 2012a. « A cognitive system for fast Quality of Transmission estimation in core optical networks ». In *Optical Fiber Communication Conference and Exposition (OFC/NFOEC), 2012 and the National Fiber Optic Engineers Conference*. (4-8 March 2012), p. 1-3.
- Jimenez, T., J. C. Aguado, I. De Miguel, R. J. Duran, N. Fernandez, M. Angelou, D. Sanchez, N. Merayo, P. Fernandez, N. Atallah, R. Lorenzo, I. Tomkos et E. J. Abril. 2012b. « Enhancing optical networks with cognition: Case-Based Reasoning to estimate the quality of transmission ». In *Cognitive Methods in Situation Awareness and Decision Support (CogSIMA), 2012 IEEE International Multi-Disciplinary Conference on*. (6-8 March 2012), p. 166-169.
- Keiser, Gerd. 2010. *Optical fiber communications* (2010), 4th ed. New York: McGraw-Hill, xxviii, 654 p. p.
- Khan, F. N., Q. Fan, C. Lu et A. P. T. Lau. 2018. « Machine Learning-Assisted Optical Performance Monitoring in Fiber-Optic Networks ». In *2018 IEEE Photonics Society Summer Topical Meeting Series (SUM)*. (9-11 July 2018), p. 53-54.
- Khan, F. N., Q. Fan, C. Lu et A. P. T. Lau. 2019. « An Optical Communication's Perspective on Machine Learning and Its Applications ». *Journal of Lightwave Technology*, vol. 37, n° 2, p. 493-516.
- Khan, Faisal Nadeem, Qirui Fan, Alan Pak Tao Lau et Chao Lu (PWO). 2020. *Applications of machine-learning in optical communications and networks*, 11309. Coll. « SPIE OPTO ». SPIE.

- Kotsiantis, S. B. 2007. « Supervised Machine Learning: A Review of Classification Techniques ». In *Proceedings of the 2007 conference on Emerging Artificial Intelligence Applications in Computer Engineering: Real Word AI Systems with Applications in eHealth, HCI, Information Retrieval and Pervasive Technologies*. p. 3–24. IOS Press.
- L. Barletta, A. Giusti, C. Rottondi, M. Tornatore,. March 2017. *QoT estimation for unestablished lighpaths using machine learning*, . Coll. « Optical Fiber Communications Conferece (OFC) 2017 ».
- Liu, Xiaomin, Huazhi Lun, Mengfan Fu, Yunyun Fan, Lilin Yi, Weisheng Hu et Qunbi Zhuge. 2020. « AI-Based Modeling and Monitoring Techniques for Future Intelligent Elastic Optical Networks ». *Applied Sciences*, vol. 10, n° 1, p. 363.
- Mahmoud, Qusay H. 2007. *Cognitive networks : Towards Self-Aware Networks*.
- Marino, R. Di, C. Rottondi, A. Giusti et A. Bianco. 2020. « Assessment of Domain Adaptation Approaches for QoT Estimation in Optical Networks ». In *2020 Optical Fiber Communications Conference and Exhibition (OFC)*. (8-12 March 2020), p. 1-3.
- Marsland, Stephen. 2014. *Machine Learning : An Algorithmic Perspective, Second Edition*. Bosa Roca, UNITED STATES: CRC Press LLC.
- Mata, J., I. d. Miguel, R. J. Durán, J. C. Aguado, N. Merayo, L. Ruiz, P. Fernández, R. M. Lorenzo, E. J. Abril et I. Tomkos. 2018a. « Supervised Machine Learning Techniques for Quality of Transmission Assessment in Optical Networks ». In *2018 20th International Conference on Transparent Optical Networks (ICTON)*. (1-5 July 2018), p. 1-4.
- Mata, J., I. de Miguel, R. J. Durán, J. C. Aguado, N. Merayo, L. Ruiz, P. Fernández, R. M. Lorenzo et E. J. Abril. 2017. « A SVM approach for lightpath QoT estimation in optical transport networks ». In *2017 IEEE International Conference on Big Data (Big Data)*. (11-14 Dec. 2017), p. 4795-4797.
- Mata, Javier, Ignacio de Miguel, Ramón J. Durán, Noemí Merayo, Sandeep Kumar Singh, Admela Jukan et Mohit Chamania. 2018b. « Artificial intelligence (AI) methods in optical networks: A comprehensive survey ». *Optical Switching and Networking*, vol. 28, p. 43-57.
- Mezni, Ameni, Douglas W. Charlton, Christine Tremblay et Christian Desrosiers. 2020. « Deep Learning for Multi-Step Performance Prediction in Operational Optical Networks

». In *Conference on Lasers and Electro-Optics*. (Washington, DC, 2020/05/10), p. STh4M.1. Coll. « OSA Technical Digest ».
 < http://www.osapublishing.org/abstract.cfm?URI=CLEO_SI-2020-STh4M.1 >.

Michel, A. 2016. « Monitoring of performance parameters in a metro ring using a 100G coherent transponder ». Final Report. École de Technologie Supérieure (ÉTS).

Mo, W., C. L. Gutterman, Y. Li, G. Zussman et D. C. Kilper. 2018a. « Deep Neural Network Based Dynamic Resource Reallocation of BBU Pools in 5G C-RAN ROADM Networks ». In *2018 Optical Fiber Communications Conference and Exposition (OFC)*. (11-15 March 2018), p. 1-3.

Mo, W., Y. Huang, S. Zhang, E. Ip, D. C. Kilper, Y. Aono et T. Tajima. 2018b. « ANN-Based Transfer Learning for QoT Prediction in Real-Time Mixed Line-Rate Systems ». In *2018 Optical Fiber Communications Conference and Exposition (OFC)*. (11-15 March 2018), p. 1-3.

Mohri, Mehryar, Afshin Rostamizadeh et Ameet Talwalkar. 2012. *Foundations of Machine Learning*. The MIT Press.

Moore, D.S., W. Notz et M.A. Fligner. 2013. *The Basic Practice of Statistics*. W.H. Freeman and Company.

Morais, R. M., et J. Pedro. 2018. « Machine learning models for estimating quality of transmission in DWDM networks ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 10, n° 10, p. D84-D99.

Mukherjee, B. 2020. *Springer Handbook of Optical Networks*. Springer.

Mukherjee, Biswanath. 2006. *Optical WDM Networks* (2006). Boston, MA: Springer Science+Business Media, Inc.

Musumeci, F., C. Rottondi, G. Corani, S. Shahkarami, F. Cugini et M. Tornatore. 2019a. « A Tutorial on Machine Learning for Failure Management in Optical Networks ». *Journal of Lightwave Technology*, vol. 37, n° 16, p. 4125-4139.

Musumeci, F., C. Rottondi, A. Nag, I. Macaluso, D. Zibar, M. Ruffini et M. Tornatore. 2019b. « An Overview on Application of Machine Learning Techniques in Optical Networks ». *IEEE Communications Surveys & Tutorials*, vol. 21, n° 2, p. 1383-1408.

Panayiotou, T., G. Savva, B. Shariati, I. Tomkos et G. Ellinas. 2019a. « Machine Learning for QoT Estimation of Unseen Optical Network States ». In *2019 Optical Fiber Communications Conference and Exhibition (OFC)*. (3-7 March 2019), p. 1-3.

- Panayiotou, T., G. Savvas, I. Tomkos et G. Ellinas. 2019b. « Centralized and Distributed Machine Learning-Based QoT Estimation for Sliceable Optical Networks ». In *2019 IEEE Global Communications Conference (GLOBECOM)*. (9-13 Dec. 2019), p. 1-7.
- Pointurier, Yvan. 2017. « Design of Low-Margin Optical Networks ». *Journal of Optical Communications and Networking*, vol. 9, n° 1, p. A9-A17.
- Proietti, R., X. Chen, K. Zhang, G. Liu, M. Shamsabardeh, A. Castro, L. Velasco, Z. Zhu et S. J. Ben Yoo. 2019. « Experimental demonstration of machine-learning-aided QoT estimation in multi-domain elastic optical networks with alien wavelengths ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 11, n° 1, p. A1-A10.
- Ratner, Bruce. 2017. *Statistical and Machine-Learning Data Mining, Third Edition: Techniques for Better Predictive Modeling and Analysis of Big Data, Third Edition*. Chapman & Hall/CRC.
- Roberts, K., Q. Zhuge, I. Monga, S. Gareau et C. Laperle. 2017. « Beyond 100 Gb/s: capacity, flexibility, and network optimization [Invited] ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 9, n° 4, p. C12-C23.
- Ruiz, M., F. Fresi, A. P. Vela, G. Meloni, N. Sambo, F. Cugini, L. Poti, L. Velasco et P. Castoldi. 2016. « Service-triggered Failure Identification/Localization Through Monitoring of Multiple Parameters ». In *ECOC 2016; 42nd European Conference on Optical Communication*. (18-22 Sept. 2016), p. 1-3.
- Samadi, P., D. Amar, C. Lepers, M. Lourdiane et K. Bergman. 2017. « Quality of Transmission Prediction with Machine Learning for Dynamic Operation of Optical WDM Networks ». In *2017 European Conference on Optical Communication (ECOC)*. (17-21 Sept. 2017), p. 1-3.
- Sambo, N., P. Castoldi, A. D' Errico, E. Riccardi, A. Pagano, M. S. Moreolo, J. M. Fàbrega, D. Rafique, A. Napoli, S. Frigerio, E. H. Salas, G. Zervas, M. Nolle, J. K. Fischer, A. Lord et J. P. F. Giménez. 2015. « Next generation sliceable bandwidth variable transponders ». *IEEE Communications Magazine*, vol. 53, n° 2, p. 163-171.
- Sartzetakis, I., K. K. Christodoulopoulos et E. M. Varvarigos. 2019. « Accurate quality of transmission estimation with machine learning ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 11, n° 3, p. 140-150.
- Sartzetakis, I., K. Christodoulopoulos, C. P. Tsekrekos, D. Syvridis et E. Varvarigos. 2016. « Quality of transmission estimation in WDM and elastic optical networks accounting

for space–spectrum dependencies ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 8, n° 9, p. 676-688.

Shahkarami, Shahin, Francesco Musumeci, Filippo Cugini et Massimo Tornatore. 2018. « Machine-Learning-Based Soft-Failure Detection and Identification in Optical Networks ». *Journal of Optical Communications and Networking*.

Shmilovici, Armin. 2005. « Support Vector Machines ». In *Data Mining and Knowledge Discovery Handbook*, sous la dir. de Maimon, Oded, et Lior Rokach. Shmilovici2005. p. 257-276. Boston, MA: Springer US. < https://doi.org/10.1007/0-387-25465-X_12 >.

Simmons, Jane M. 2014. *Optical Network Design and Planning* (2014), 2nd ed. 2014. Cham: Springer International Publishing, (XXV, 516 p.) p.

Singh, A., N. Thakur et A. Sharma. 2016. « A review of supervised machine learning algorithms ». In *2016 3rd International Conference on Computing for Sustainable Global Development (INDIACom)*. (16-18 March 2016), p. 1310-1315.

Stańczyk, Urszula, et L. C. Jain. 2015. *Feature selection for data and pattern recognition*. Heidelberg: Springer. In WorldCat.org.

T. Panayiotou, S. P. Chatzis, G. Ellinas. 2017. « Performance Analysis of a Data-Driven Quality-of-Transmission Decision Approach on a Dynamic MulticastCapable Metro Optical Network ». *Journal of Optical Communications and Networking*, vol. 9, n° 1.

The Weather Company. 2014.
< <https://www.wunderground.com/history/monthly/ca/montreal/CYUL/date> >.

Thomas, Ryan, Daniel Friend, Luiz Dasilva et Allen Mackenzie. 2006. « Cognitive networks: adaptation and learning to achieve end-to-end performance objectives ». *IEEE Communications Magazine*, vol. 44, n° 12, p. 51-57.

Tremblay, C., A. Cassidy, A. Mortelette, C. Amelin, M. Lyonnais, T. Tam et M. P. Bélanger. November 7-10, 2012. « Advanced performance measurements using 40G coherent systems ». In *IEEE/OSA/SPIE Asia Communications and Photonics (ACP'12) Conference*. (Guangzhou, China). Vol. Paper AS3A.1.

Tremblay, C., A. Michel, M. J. Tanoh, M. P. Bélanger, S. Clarke, D. W. Charlton, D. L. Peterson et G. A. Wellbrock. 2017. « Dynamics of polarization fluctuations in aerial and buried links ». In *2017 19th International Conference on Transparent Optical Networks (ICTON)*. (2-6 July 2017), p. 1-1.

- Tremblay, C., A. Michel, M. J. Tanoh, M. P. Bélanger, S. Clarke, D. W. Charlton, D. L. Peterson et G. A. Wellbrock. July 3-6, 2017. « Dynamics of Polarization Fluctuations in Aerial and Buried Links ». In *IEEE Int. Conf. Transparent Optical Networks (ICTON'17)*. (Girona, Spain).
- Tremblay, Christine, Stéphanie Allogba et Sandra Aladin. 2019. « Quality of Transmission Estimation and Performance Prediction of Lightpaths Using Machine Learning ». In *45th European Conf. Optical Commun. (ECOC)*. (Dublin, Ireland).
- Vejdannik, Masoud, et Ali Sadr. 2020. « Modular neural networks for quality of transmission prediction in low-margin optical networks ». *Journal of Intelligent Manufacturing*.
- Vela, A. P., M. Ruiz, F. Cugini et L. Velasco. 2017a. « Combining a machine learning and optimization for early pre-FEC BER degradation to meet committed QoS ». In *2017 19th International Conference on Transparent Optical Networks (ICTON)*. (2-6 July 2017), p. 1-4.
- Vela, A. P., M. Ruiz, F. Fresi, N. Sambo, F. Cugini, G. Meloni, L. Potì, L. Velasco et P. Castoldi. 2017b. « BER Degradation Detection and Failure Identification in Elastic Optical Networks ». *Journal of Lightwave Technology*, vol. 35, n° 21, p. 4595-4604.
- Vela, Alba P., Marc Ruiz, Francesco Fresi, Nicola Sambo, Filippo Cugini, Luis Velasco et Piero Castoldi. 2017c. « Early Pre-FEC BER Degradation Detection to Meet Committed QoS ». In *Optical Fiber Communication Conference*. (Los Angeles, California, 2017/03/19), p. W4F.3. Coll. « OSA Technical Digest (online) »: Optical Society of America. < <http://www.osapublishing.org/abstract.cfm?URI=OFC-2017-W4F.3> >.
- Velasco, L., A. C. Piat, O. Gonzalez, A. Lord, A. Napoli, P. Layec, D. Rafique, A. D'Errico, D. King, M. Ruiz, F. Cugini et R. Casellas. 2019. « Monitoring and Data Analytics for Optical Networking: Benefits, Architectures, and Use Cases ». *IEEE Network*, vol. 33, n° 6, p. 100-108.
- Vinutha, H. P., B. Poornima et B. M. Sagar. 2018. « Detection of Outliers Using Interquartile Range Technique from Intrusion Dataset ». In (Singapore). 10.1007/978-981-10-7563-6_53. p. 511-518. Coll. « Information and Decision Sciences »: Springer Singapore.
- Wang, Y., J. Zhang, H. Zhu, M. Long, J. Wang et P. S. Yu. 2019. « Memory in Memory: A Predictive Neural Network for Learning Higher-Order Non-Stationarity From Spatiotemporal Dynamics ». In *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. (15-20 June 2019), p. 9146-9154.

- Wang, Zhilong, Min Zhang, Danshi Wang, Chuang Song, Min Liu, Jin Li, Liqi Lou et Zhuo Liu. 2017. « Failure prediction using machine learning and time series in optical network ». *Optics Express*, vol. 25, n° 16, p. 18553-18565.
- Weiss, Karl, Taghi M. Khoshgoftaar et DingDing Wang. 2016. « A survey of transfer learning ». *Journal of Big Data*, vol. 3, n° 1, p. 9.
- Witten, Ian H., Eibe Frank et Mark A. Hall (665). 2011. *Data Mining: Practical Machine Learning Tools and Techniques - 3rd ed.*: Morgan Kaufmann.
- Wu, Jia, Xiu-Yun Chen, Hao Zhang, Li-Dong Xiong, Hang Lei et Si-Hao Deng. 2019. « Hyperparameter Optimization for Machine Learning Models Based on Bayesian Optimizationb ». *Journal of Electronic Science and Technology*, vol. 17, n° 1, p. 26-40.
- Yaméogo, B. L. M., D. W. Charlton, D. Doucet, C. Desrosiers, M. O' Sullivan et C. Tremblay. 2020. « Trends in Optical Span Loss Detected Using the Time Series Decomposition Method ». *Journal of Lightwave Technology*, vol. 38, n° 18, p. 5026-5035.
- Yu, J., W. Mo, Y. Huang, E. Ip et D. C. Kilper. 2019. « Model transfer of QoT prediction in optical networks based on artificial neural networks ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 11, n° 10, p. C48-C57.
- Zhang, C., D. Wang, L. Wang, J. Song, S. Liu, J. Li, L. Guan, Z. Liu et M. Zhang. 2020. « Temporal data-driven failure prognostics using BiGRU for optical networks ». *IEEE/OSA Journal of Optical Communications and Networking*, vol. 12, n° 8, p. 277-287.
- Zhuang, F., Z. Qi, K. Duan, D. Xi, Y. Zhu, H. Zhu, H. Xiong et Q. He. 2021. « A Comprehensive Survey on Transfer Learning ». *Proceedings of the IEEE*, vol. 109, n° 1, p. 43-76.

