



Estimation et selection pour les modeles additifset application a la prevision de la consommation electrique

Vincent Thouvenot

► To cite this version:

Vincent Thouvenot. Estimation et selection pour les modeles additifset application a la prevision de la consommation electrique. Statistiques [stat]. Paris Saclay, 2015. Français. NNT: . tel-01708576

HAL Id: tel-01708576

<https://inria.hal.science/tel-01708576>

Submitted on 13 Feb 2018

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

THÈSE

Présentée pour obtenir

LE GRADE DE DOCTEUR EN SCIENCES DE
L'UNIVERSITÉ PARIS-SUD XI

Spécialité: Mathématiques

par

Vincent THOUVENOT

Estimation et sélection pour les modèles additifs et application à la prévision de la consommation électrique

Soutenue le 17 Décembre 2015 devant la Commission d'examen:

M.	Anestis ANTONIADIS	(Directeur de thèse)
Mme.	Irène GIJBELS	(Rapporteur)
M.	Yannig GOUDE	(Responsable industriel)
M.	Pascal MASSART	(Président du jury)
M.	Eric MATZNER-LOBER	(Examineur)
M.	Pierre PINSON	(Rapporteur)
M.	Jean-Michel POGGI	(Directeur de thèse)

Rapporteurs:

Mme.	Irène Gijbels
M.	Pierre Pinson



Thèse préparée au
Département de Mathématiques d'Orsay
Laboratoire de Mathématiques (UMR 8628), Bât. 425
Université Paris-Sud 11
91 405 Orsay CEDEX

Résumé

L'électricité ne se stockant pas aisément, EDF a besoin d'outils de prévision de consommation et de production efficaces. Le développement de nouvelles méthodes automatiques de sélection et d'estimation de modèles de prévision est nécessaire. En effet, grâce au développement de nouvelles technologies, EDF peut étudier les mailles locales du réseau électrique, ce qui amène à un nombre important de séries chronologiques à étudier. De plus, avec les changements d'habitude de consommation et la crise économique, la consommation électrique en France évolue. Pour cette prévision, nous adoptons ici une méthode semi-paramétrique à base de modèles additifs. L'objectif de ce travail est de présenter des procédures automatiques de sélection et d'estimation de composantes d'un modèle additif avec des estimateurs en plusieurs étapes. Nous utilisons du Group LASSO, qui est, sous certaines conditions, consistant en sélection, et des P-Splines, qui sont consistantes en estimation. Nos résultats théoriques de consistance en sélection et en estimation sont obtenus sans nécessiter l'hypothèse classique que les normes des composantes non nulles du modèles additifs soient bornées par une constante non nulle. En effet, nous autorisons cette norme à pouvoir converger vers 0 à une certaine vitesse. Les procédures sont illustrées sur des applications pratiques de prévision de consommation électrique nationale et locale.

Mots-clés. Group LASSO, Estimateurs en plusieurs étapes, Modèle Additif, Prévision de charge électrique, P-Splines, Sélection de variables

Abstract

French electricity load forecasting encounters major changes since the past decade. These changes are, among others things, due to the opening of electricity market (and economical crisis), which asks development of new automatic time adaptive prediction methods. The advent of innovating technologies also needs the development of some automatic methods, because we have to study thousands or tens of thousands time series. We adopt for time prediction a semi-parametric approach based on additive models. We present an automatic procedure for covariate selection in a additive model. We combine Group LASSO, which is selection consistent, with P-Splines, which are estimation consistent. Our estimation and model selection results are valid without assuming that the norm of each of the true non-zero components is bounded away from zero and need only that the norms of non-zero components converge to zero at a certain rate. Real applications on local and aggregate load forecasting are provided.

Keywords. Additive Model, Group LASSO, Load Forecasting, Multi-stage estimator, P-Splines, Variables selection

Remerciements

Au cours de mes trois ans de travaux de thèse, j'ai pu compter sur la présence de nombreuses personnes, que je tiens à remercier dans ce paragraphe.

En premier lieu, je tiens à exprimer toute ma plus profonde gratitude à Anestis Antoniadis et Jean-Michel Poggi, mes deux directeurs de thèse, ainsi qu'à Yannig Goude, mon référent industriel. J'ai eu la chance d'avoir un encadrement de très haute qualité. Tout au long de ma thèse, vous m'avez, chacun à votre manière, soutenu et éclairé de très nombreux conseils et idées. Merci de votre patience et de la très forte implication dont vous avez fait preuve tout au long de la thèse, ainsi que pour votre disponibilité et les très nombreuses heures que vous m'avez consacré. Je vous remercie de m'avoir cadré quand je m'éparpillais et permis d'aller à des conférences très intéressantes. Travailler avec vous m'a énormément apporté, tant professionnellement que humainement. J'espère que j'aurais encore l'occasion de travailler avec vous après cette thèse.

Je tiens aussi à remercier chaleureusement Pascal Massart, qui m'a fait l'honneur de présider mon jury de thèse, ainsi que Pierre Pinson et Irène Gijbels, qui ont accepté de rapporter ma thèse. Mes remerciements vont aussi à Eric Matzner-Lober qui a accepté d'être membre de mon jury de thèse. Je remercie également ce dernier d'avoir accepté que je fasse une présentation dans la session qu'il organisait à la conférence ISNPS à Cadix.

Je remercie sincèrement l'ensemble des membres du groupe de prévision de consommation court et moyen terme d'OSIRIS à EDF R&D, qu'ils soient agents, stagiaires ou prestataires. J'ai eu l'occasion de beaucoup apprendre au contact de chacun d'entre vous. Je remercie également les agents pour leur accueil bienveillant. Je tiens tout particulièrement à exprimer ma vive reconnaissance à Audrey, l'assistante du groupe, qui m'a sauvé des méandres administratifs d'EDF à de très nombreuses reprises, et toujours avec la bonne humeur, la patience et la grande efficacité qui la caractérisent. Merci à Pierre avec qui j'ai eu la chance de partager deux bureaux (le froid de notre premier bureau t'a bien préparé pour ta future destination!) et auprès de qui j'ai eu l'occasion de beaucoup apprendre. Merci également à (l'autre) Audrey, avec qui j'ai eu le plaisir de travailler ponctuellement, profitant à chaque fois de ses connaissances et de sa disponibilité. Je remercie aussi Titouan pour l'intervention efficace réalisée sur le code R en toute fin de thèse. Merci à Xavier pour les quelques discussions que nous avons pu avoir, ainsi que pour les conseils avisés. Merci à Cyrille et à Anne pour le temps passé ensemble lors des conférences respectivement à Aalborg et à Lille. Merci à mes trois co-bureaux de stage du B314 pour les bons moments passés ensemble durant cette période. Merci enfin Raphaël, Amandine, Yohann, Julien et les autres pour les quelques verres partagés après le travail.

J'ai passé moins de temps à Orsay qu'à EDF durant ma thèse. Cependant, je voudrais remercier Elodie, Emilie, Vincent, Lucie, Clément, Alba et Thomas, avec qui j'ai toujours eu plaisir à discuter. Merci aussi à Valérie qui a toujours répondu efficacement à tous mes problèmes administratifs à Orsay.

Pendant ces deux dernières années, j'ai eu l'occasion d'enseigner à l'ENSAI. Merci Myriam de m'avoir fait confiance et de m'en avoir donné l'opportunité. Ce fut une expérience particulièrement enrichissante et enthousiasmante. Je te remercie aussi pour ta disponibilité et ton efficacité à l'époque où j'étudiais à l'ENSAI. J'en profite également pour remercier Lise, aller donner les TDs avec toi a été un vrai plaisir et je te remercie pour nos déjeuners à Clamart ainsi que pour ton soutien et tes nombreux conseils. Je remercie enfin les professeurs de l'ENSAI qui m'ont donné le goût des Statistiques, et plus généralement les enseignants qui m'ont transmis leurs connaissances.

Évidemment, ces remerciements seraient incomplets si je ne remerciais pas mes proches, dont le soutien affectif et la présence sont un moteur essentiel pour moi. Merci donc à mes amis qui

ont eu un rôle fondamental durant ma thèse, et notamment à Nicolas et Claire, votre aide pour la relecture de la thèse, ainsi que votre présence et votre soutien permanent durant les trois dernières années, et plus particulièrement au cours de ces derniers mois, m'ont énormément apporté. Merci également Alice pour tes encouragements, tes conseils et ton écoute. Je ne peux malheureusement pas faire une liste exhaustive, mais merci également à Gaetan, Laura, Laurent, Benoît et plus généralement aux autres coureurs et joueurs de tennis, ainsi qu'aux personnes présentes lors des voyages à Lisbonne et à Londres, de m'avoir permis de couper et faire des pauses. Également, je profite de ce paragraphe pour adresser toute ma reconnaissance à ma famille, ma mère, mon père, mon frère et ma soeur, ainsi qu'à ses deux enfants, pour le soutien permanent et inconditionnel. J'exprime enfin une pensée pour ma grand mère.

Publications and main activities

Article to appear

V. Thouvenot, A. Pichavant, Y. Goude, A. Antoniadis, and J-M Poggi. Electricity forecasting using multi-stage estimators of nonlinear additive models. *To appear in IEEE Transactions on Power Systems*, 2015

Article submitted in proceedings of an international conference

A. Antoniadis, X. Brossat, Y. Goude, J-M. Poggi, and V. Thouvenot. Automatic component selection in additive modeling of french national electricity load forecasting. In *Proceedings of ISNPS 2014*, 2015

Technical report

A. Antoniadis, Y. Goude, J-M. Poggi, and V. Thouvenot. Sélection de variables dans les modèles additifs avec des estimateurs en plusieurs étapes. <https://hal.archives-ouvertes.fr/hal-01116100>

Conference presentations

A. Antoniadis, X. Brossat, Y. Goude, J-M. Poggi, and V. Thouvenot. *R package CASA : Component Automatic Selection in Additive models*, UseR! à Aalborg, 2015

A. Antoniadis, X. Brossat, Y. Goude, J-M. Poggi, and V. Thouvenot. *Estimateurs en plusieurs étapes de modèles additifs appliqués à la modélisation et à la prévision de la consommation électrique*, 47ème JdS à Lille, 2015

A. Antoniadis, X. Brossat, Y. Goude, J-M. Poggi, and V. Thouvenot. *Automatic component selection in additive modeling of French national electricity load forecasting*, ISNPS à Cadix, 2014

A. Antoniadis, X. Brossat, Y. Goude, J-M. Poggi, and V. Thouvenot. *Sélection automatique de composantes dans les modèles additifs : application à la consommation nationale d'électricité*, 46ème JdS à Rennes, 2014

Prototype of R package

CASA : Component Automatic Selection in Additive models.



Table des matières

Introduction	11
0.1 Contexte	11
0.1.1 Contexte industriel	11
0.1.2 Etat de l'art sur la prévision de consommation à EDF R&D	14
0.2 Cadre statistique, estimateurs en plusieurs étapes et illustration sur des simulations	17
0.2.1 Modèle additif et méthodes de régression	17
0.2.2 Estimateurs en plusieurs étapes	18
0.2.3 Perspectives méthodologiques	20
0.3 Estimateurs en plusieurs étapes de modèles additifs appliqués à la prévision de consommation électrique à plusieurs niveaux d'agrégation	20
0.3.1 Application agrégée au niveau national	21
0.3.2 Application locale	21
0.3.3 Résumé des résultats et implémentation des méthodes	23
0.3.4 Autres applications	24
0.4 Un résultat de consistance pour un estimateur en plusieurs étapes	24
1 Cadre statistique, estimateurs en plusieurs étapes et illustration sur des simulations	27
Introduction	27
1.1 Généralités sur le modèle linéaire	27
1.2 Généralités sur le modèle additif	36
1.3 Procédures de sélection et d'estimation	41
1.4 Simulations	51
2 Estimateurs en plusieurs étapes de modèles additifs appliqués à la prévision de consommation à plusieurs niveaux d'agrégation	63
Introduction	63
2.1 Modélisation et prévision de la consommation pour le portefeuille EDF	64
2.1.1 Description des données	64
2.1.2 Modèle paramétrique historique d'EDF	67
2.1.3 Modèle additif moyen terme d'EDF construit grâce à l'expertise métier . .	69
2.1.4 Modèles additifs moyen terme construits avec une sélection automatisée . .	70
2.1.5 Modèles court terme	78

2.2	Modélisation et prévision de consommation pour des mailles locales de réseaux de distribution	82
2.2.1	Description des données	83
2.2.2	Les modèles de référence	85
2.2.3	Application des procédures de sélection automatique sur GEFCom 2012 . .	86
2.2.4	Application sur les postes sources	92
	Conclusion	94
3	Un résultat de consistance pour un estimateur en plusieurs étapes	97
3.1	Introduction	97
3.2	Etude asymptotique des procédures	101
3.2.1	Hypothèses pour l'étude asymptotique	102
3.2.2	Consistance asymptotique en sélection du BIC appliqué à l'estimateur des MCO	104
3.2.3	Etude de la procédure Post1	109
3.3	Conclusion	113
3.4	Démonstration pour la consistance asymptotique en sélection du BIC appliqué à l'EMCO	114
3.4.1	Cas où $S^* \not\subseteq S$	114
3.4.2	Cas où $S^* \subseteq S$	121
3.5	Démonstration de la consistance en sélection de la procédure Post1	123
	Perspectives	131
A	Annexe : Vignette R CASA	135
	Introduction	135
A.1	Statistical framework and corresponding R packages	135
A.1.1	Additive models	135
A.1.2	Multi-stage estimator : the Post2 procedure	136
A.2	Overview of the R package CASA	137
A.2.1	A real data example of CASA 's use	138
A.2.2	Two main functions	147
A.2.3	Questions and Answers	149
A.2.4	Perspectives	149
A.3	<i>Post2</i> help	149
B	Annexe : Résultats complets des simulations présentées dans le chapitre 1	151
C	Annexe : Résultats complémentaires du chapitre 2	155
C.1	Lexique	155
C.2	Evaluation des performances en prévision des modèles	156
C.3	Portefeuille EDF	157
C.4	GEFCom 2012	163

D Annexe : Simulation associée au chapitre 3	167
Références	169



Introduction

Le travail présenté ici a été réalisé dans le cadre d'une thèse CIFRE à EDF R&D à Clamart et au Laboratoire de Mathématiques de l'Université d'Orsay. Il porte sur la sélection et l'estimation de composantes non nulles de modèles additifs creux à l'aide d'estimateurs en plusieurs étapes. Les méthodes ont été proposées, analysées théoriquement, expérimentées sur des simulations et appliquées à différents jeux de données de consommation électrique internes et externes à EDF. Dans cette introduction générale, nous présentons les contextes industriel et statistique de la thèse puis nous présentons les trois chapitres qui la composent. Dans le premier chapitre, nous posons le cadre statistique dans lequel nous nous plaçons et développons les méthodes étudiées ensuite. Nous testons ces méthodes sur des simulations. Dans le chapitre 2, nous nous concentrons sur une sous-famille d'estimateurs en plusieurs étapes et l'utilisons sur différents jeux de données de prévision de consommation électrique. Enfin, dans le dernier chapitre, nous établissons des propriétés de consistance d'un estimateur en plusieurs étapes.

0.1 Contexte

0.1.1 Contexte industriel

La prévision de consommation électrique est une activité cruciale pour une entreprise comme EDF. Elle est nécessaire pour des objectifs aussi variés que le management de la production, la gestion et la maintenance du réseau électrique, le marché de l'électricité et la tarification. Il existe peu de solutions de stockage de l'électricité (barrage hydraulique, batterie, voir [27]) et celles-ci sont généralement insuffisantes pour répondre à la demande totale de l'électricité. L'équilibre offre-demande doit donc toujours être garanti afin d'éviter les risques physiques (black-out partiel ou total) ou financiers (pénalités).

Cycles et variables influençant la consommation La consommation d'électricité française varie en fonction de plusieurs phénomènes et selon les trois cycles temporels suivants :

- cycle annuel, avec une pointe de consommation en janvier et un creux autour du 15 août. Ce cycle s'explique par l'activité économique (creux du 15 août par exemple) et par le climat (température plus froide en hiver) ;
- cycle hebdomadaire, avec une charge consommée plus forte les jours ouvrables que les week-ends et les jours fériés (activité économique) ;
- cycle journalier, avec une plus forte consommation le jour que la nuit avec des pics vers 9h le matin (les gens arrivent au travail), vers 19h (les gens rentrent du travail) et 22h (allumage des chauffe-eaux).

Ces trois cycles sont représentés sur la Figure 1. La partie gauche de cette figure présente la charge journalière moyenne consommée pour le portefeuille EDF (proche des données nationales) en 2008 et la partie droite, la charge consommée durant deux semaines en 2008 : une en hiver et l'autre en été.

Le climat a un impact sur la consommation nationale. La température influence l'utilisation du

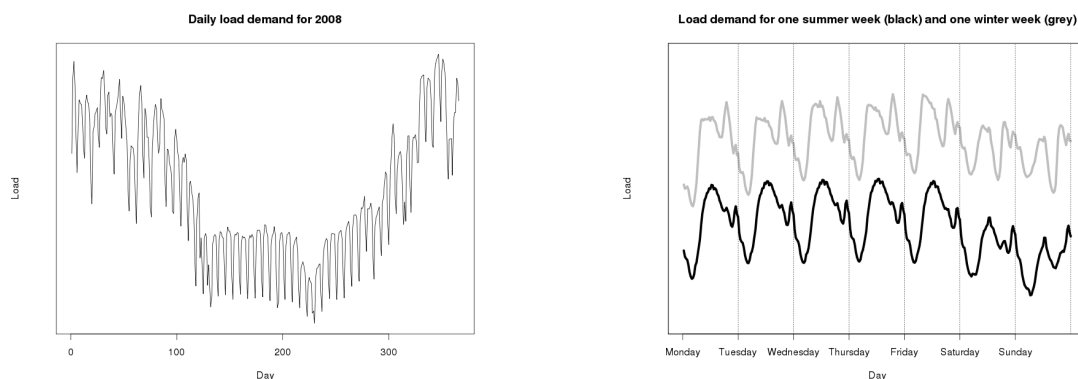


FIGURE 1 – Charge du portefeuille consommée en 2008 (à gauche) et sur deux semaines (à droite) : une pendant l’hiver (gris) et l’autre pendant l’été 2008 (noir)

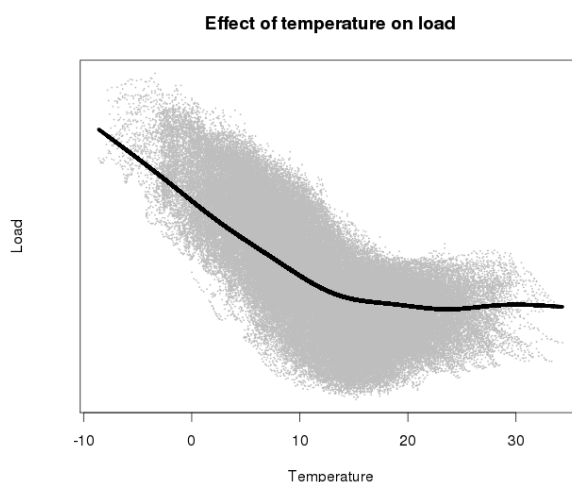


FIGURE 2 – Effet estimé de la température (en degré Celsius) sur la charge consommée en 2008 sur le portefeuille EDF

chauffage, majoritairement électrique en France, et de la climatisation. La nébulosité modifie les effets du rayonnement solaire dans les habitations et influence donc les utilisations du chauffage et de l’éclairage. La Figure 2 présente la charge consommée en fonction de la température et illustre le lien non linéaire existant entre la température et la charge électrique consommée pour le portefeuille EDF.

Naturellement, d’autres facteurs, tels que la tarification, le changement d’heure ou des événements exceptionnels (grand rendez-vous sportif ou culturel par exemple) influencent aussi la charge consommée.

Horizons de prévision Si les variables calendaires sont déterministes et connues à l’avance, les variables météorologiques sont aléatoires. Des méthodes différentes sont utilisées pour les prévoir selon la période de prévision étudiée (voir Annexe C.1). Il y a donc une notion d’horizon de prévision pour la météorologie, qui se retrouve pour la prévision de consommation électrique. Celle-ci est fortement dépendante de la situation économique, qui évolue généralement très peu du jour pour le lendemain, mais qui peut beaucoup changer en 15 ans. Les modèles ainsi que les objectifs sont différents selon la période prédite (voir par exemple le chapitre 12 de [30] ou [121]). Nous distinguons ainsi quatre horizons de prévision.

La prévision de la consommation électrique à long terme (à un horizon d'une ou de quelques dizaines d'années) est nécessaire pour les investissements futurs d'EDF. RTE (Réseau de Transport d'Électricité), filiale d'EDF gérant le réseau de transport, établit tous les deux ans un bilan prévisionnel pluriannuel de l'équilibre offre/demande d'électricité en France, assurant ainsi l'équilibre entre l'offre et la demande pour les 15 prochaines années. La prévision à long terme de la consommation est un des maillons nécessaires pour vérifier que le parc productif futur sera suffisant pour répondre à la demande future. Nous ne traitons pas de cet horizon dans ce document.

La prévision moyen terme de la consommation d'électricité consiste en une prévision à l'horizon de quelques semaines à une année. Elle est nécessaire pour la gestion et la maintenance du réseau (arrêt ou démarrage de tranche de centrales nucléaires, travaux, etc), ainsi que pour la gestion du portefeuille et du marché électrique. L'étude de cet horizon sert à quantifier les risques de défaillances physiques ainsi que l'équilibre du périmètre EDF. Les pics et les ruptures de consommation intéressant EDF, les pas d'échantillonnage des études sont fins (pas d'une heure, de 30 minutes voire de 10 minutes).

La prévision à un horizon hebdomadaire et bi-hebdomadaire permet de quantifier les risques de défaillances physiques, de modéliser le périmètre d'équilibre d'EDF, de placer des effacements (voir explication dans le chapitre 2) et de réaliser des achats sur le marché de l'électricité.

La prévision à court terme (à un horizon de 24 heures) est utilisée pour optimiser la charge à produire le lendemain et s'assurer un approvisionnement suffisant d'électricité pour répondre à la demande tout en minimisant les coûts. La prévision journalière permet d'introduire les derniers ajustements pour assurer l'équilibre du périmètre EDF, de gérer la production court terme et de réaliser les derniers achats et décisions d'effacement.

Évolutions récentes Depuis une dizaine d'années, plusieurs évolutions majeures ont eu lieu et ont complexifié la problématique de prévision de la consommation électrique. Le récent développement de nombreuses technologies de mesure permet de construire des réseaux "intelligents", avec plus d'informations de plus en plus locales, et ouvre de nouvelles perspectives pour le pilotage de l'énergie et donc pour la prévision de la consommation électrique française. L'exemple le plus marquant de cette évolution est, en France, le déploiement et l'expérimentation par ERDF (Électricité Réseau Distribution de France, qui gère le réseau de distribution) de 250 000 compteurs intelligents LINKY dans la région lyonnaise. Cette expérimentation étant concluante, les autorités publiques ont fixé pour objectif l'installation de plus de 35 millions de compteurs LINKY avant 2020. Ces compteurs offrent la possibilité d'avoir quasiment en temps réel les courbes de charge individuelle. Grâce aux compteurs intelligents, les consommateurs pourront devenir acteurs de leur consommation et les fournisseurs auront des nouvelles possibilités de tarifications dynamiques permettant de mieux gérer la pointe journalière de consommation, par exemple. Les compteurs LINKY pourront être un outil de pilotage de charges, bien que l'exemple italien, où des compteurs intelligents ont déjà été déployés à grande échelle depuis 2006 par Enel, montre que la mise au point peut être longue. Le passage à la prévision locale offre de nouvelles possibilités en permettant de gérer localement le réseau électrique, et ainsi de mieux intégrer la production locale et distribuée (énergies renouvelables). Les prévisions agrégées peuvent être améliorées en agrégeant les prévisions locales. Cependant, le passage à l'étude des mailles locales conduit à des nouvelles problématiques dont :

- À un niveau de maille donné du réseau, il peut y avoir plus 2000 (postes sources) ou de 20000 (départs des postes sources), voire plus, séries chronologiques à étudier, ce qui conduit à la nécessaire **automatisation** des méthodes.
- Arrivée massive et parfois chaotique (problèmes techniques, incertitudes sur les mesures,...) de données, ce qui pose des problèmes pour le stockage des données et sur la manière de traiter rapidement ces données, conduisant à des problèmes propres au **Big Data** (voir [18]).

- Courbes de charge où le client (notamment industriel) peut être identifié, ce qui conduit à des problèmes de **confidentialité**.
- Plus les données sont locales, plus elles risquent d'être **bruitées**.

Actuellement, le paysage de la consommation nationale d'électricité évolue rapidement. Les habitudes des consommateurs, qui changeaient peu et lentement par le passé, ont été complètement bouleversées par l'apparition de nouveaux usages tels que le développement des voitures électriques, des écrans plats, des climatiseurs, des smart-phones, de bâtiments moins sensibles aux changements des températures extérieures, de l'auto-consommation (électricité produite par le consommateur pour ses propres besoins), des ampoules basse consommation, etc. En parallèle de ces changements d'habitudes, l'ouverture du marché à la concurrence, en 2004 pour les gros consommateurs et en 2007 pour les petits consommateurs, conduit à la possibilité de gains et de pertes de clients pour le portefeuille EDF et au développement d'un marché électrique qui impose de nouvelles contraintes à EDF. La fin de certains tarifs réglementés en 2016 va conduire à un changement de périmètre du portefeuille EDF. La transition énergétique va aussi influencer la manière de consommer. Fin 2014, le Conseil de l'Europe a fixé pour objectifs de diminuer les émissions de gaz à effet de serre de 40%, d'augmenter la part de l'énergie renouvelable dans la consommation d'énergie de 27% et d'améliorer l'efficacité énergétique de 27% d'ici à 2030 par rapport à des scénarios tendanciels. Le dernier objectif va directement influencer la consommation d'électricité et plus généralement d'énergie. Avec l'augmentation de la part des énergies renouvelables intermittentes, une vision probabiliste de la consommation devient nécessaire pour tenir compte des nouveaux risques réseaux. La prise en compte des aléas climatiques est nécessaire pour bien simuler simultanément la production et la consommation.

La multiplication des nouvelles applications possibles (petits agrégats de consommateurs, stations assurant la liaison entre les réseaux de transport et de distribution, départ basse tension de ces stations, etc), les changements des habitudes des consommateurs et la sortie du monopole, conduisant à une variation du portefeuille de clients EDF, amènent un besoin nouveau de méthodes s'adaptant automatiquement aux données.

0.1.2 Etat de l'art sur la prévision de consommation à EDF R&D

La prévision de consommation, particulièrement étudiée depuis une vingtaine d'années, possède une littérature riche. Nous présentons dans ce paragraphe des méthodes classiques de prévision de consommation électrique, puis nous nous intéressons plus spécifiquement à certaines méthodes établies à EDF pour contextualiser la thèse, tant dans sa dimension recherche que développement. Comme celles-ci, la thèse s'inscrit dans le cadre de recherche de méthodes automatisées capables de facilement s'adapter aux changements d'habitudes des consommateurs et de périmètre du portefeuille de clients ainsi que permettant d'étudier de nombreuses séries chronologiques tout en limitant les interventions humaines au maximum.

Pour la prévision court terme, des méthodes classiques de séries temporelles, telles que les modèles SARIMA (voir [92] ou [68]) ou les modèles de lissage exponentiel (voir [108] ou [109]) sont utilisées. La charge consommée dépend de variables exogènes (température ou nébulosité par exemple). Des modèles de régression, utilisés dans le contexte des séries temporelles, ont été proposés pour la prévision à court et moyen termes (voir [22], [24] ou [98]). Des modèles de régression non-paramétrique ont montré de bonnes performances. Par exemple, [95] présente des estimateurs à noyaux d'un modèle autorégressif non linéaire utilisé pour prévoir la consommation française à court terme ou [46] et [69] utilisent des modèles additifs pour prévoir la consommation électrique en Australie respectivement court et long termes. Récemment, des méthodes d'Intelligence Artificielle ont été utilisées : les SVM par [26], les réseaux de neurones par [45] et [73], les algorithmes évolutionnaires par [60].

La liste n'est évidemment pas exhaustive. Le lecteur peut se référer à [56] qui propose une bibliographie plus complète des méthodes de prévision de consommation électrique.

Le modèle de régression de prévision de consommation électrique historiquement utilisé à EDF est présenté par Bruhns *et al.* [21] (2005). Ces performances sont très bonnes sur les données de consommation française et sur le portefeuille EDF, mais étant paramétrique, il est peu évolutif. Pour répondre aux nouveaux défis industriels, EDF a développé des méthodes publiées dans des articles et/ou des thèses CIFRE. Nous présentons d'abord des méthodes appliquées à la prévision de consommation court terme d'un haut niveau d'agrégation (données nationales ou du portefeuille EDF), puis deux méthodes appliquées à l'étude de la consommation locale. Enfin, nous présentons des cas d'applications de modèles additifs tant à la prévision court et moyen termes des consommations électriques locales et nationale. L'ensemble des méthodes présentées s'inscrit dans le cadre de la recherche de méthodes capables de rapidement s'adapter aux changements actuels de la consommation électrique française.

0.1.2.1 Prévision nationale court terme

Modèle de séries temporelles Pour s'affranchir des problèmes du fléau de la dimension des méthodes non-paramétriques, Lefieux étudie dans sa thèse CIFRE [82] (2007), réalisée à RTE et à l'Université de Rennes, la méthode semi-paramétrique MAVE, fondée sur la notion de "directions révélatrices". Cette méthode n'obtenant pas des résultats satisfaisants en présence de variables exogènes, Lefieux propose un modèle semi-linéaire à directions révélatrices multiples. Cette thèse s'inscrit dans le cadre d'ajout de variables exogènes (vent, nébulosité,...) et de recherche de modèles plus évolutifs que le modèle historique présenté par [21].

Modèle dynamique Pour lisser les évolutions de la courbe de charge nationale dues aux changements des habitudes des consommateurs, la thèse CIFRE de Dordonnat [37] (2009), réalisée à EDF et à l'Université d'Amsterdam, utilise des modèles de régression dynamiques périodiques linéaires et non linéaires (voir aussi Dordonnat *et al.* [38], 2008). Même si cette thèse a été appliquée sur les données nationales et non sur le portefeuille EDF, les méthodes proposées s'inscrivent parfaitement dans le cadre de l'ouverture du marché électrique et sont efficaces pendant les périodes de l'année les plus variables (vacances d'été), ainsi que pour les effets variant peu à long terme (effets climatiques).

Modèles fonctionnels Les méthodes fonctionnelles permettent de transformer des modèles dynamiques non linéaires en modèles dynamiques linéaires grâce à des projections judicieuses.

En s'inspirant d'Antoniadis *et al.* [11] (2006), qui travaillent notamment sur la consommation d'électricité court terme de Paris, Cho *et al.* [29] (2012) et Cho *et al.* [28] (2015) utilisent des méthodes de Curve Linear Regression (CLR) et modélisent la charge consommée à court terme en France par un processus fonctionnel en cherchant des similitudes entre les courbes de charge journalière après décomposition dans des bases d'ondelettes. Cette méthode a des capacités prédictives concurrentielles par rapport au modèle opérationnel utilisé à EDF, sans nécessiter d'expertise métier, et s'adapte efficacement aux changements.

L'intégration dans les modèles de prévision de consommation électrique des covariables, notamment météorologiques, est généralement problématique. Dans la thèse CIFRE de Cugliari [32] (2011), réalisée à EDF et à l'Université d'Orsay, ainsi que dans Antoniadis *et al.* [3] (2012) [5] (2014), un modèle de prévision (KWF) pour séries chronologiques fonctionnelles en présence de non stationnarités, ne demandant pas l'introduction de variables exogènes pour être compétitif, est utilisé sur les données françaises et du portefeuille EDF de consommation électrique. La consommation est décomposée dans des bases d'ondelettes. La présence de groupes de données considérés comme des classes de stationnarité (e.g. le dimanche) explique en partie les non stationnarités des

séries. Cugliari propose des algorithmes de classification non supervisée pour trouver ces classes (voir [4], 2013) en utilisant une distance entre les décompositions dans les bases d'ondelettes de la consommation.

Modèle bayésien Avec la fin du monopole, le périmètre du portefeuille EDF évolue. La thèse CIFRE de Launay [77] (2012), réalisée à EDF et à l'Université de Nantes, s'inscrit bien dans ce changement de périmètre, puisqu'elle permet d'améliorer les prévisions en situation d'historique court (voir aussi [78], 2012). Launay travaille aussi sur la réalisation de prévisions en ligne à l'aide de filtres particuliers [79] (2012), développant ainsi une méthode adaptative. Le travail de Launay s'appuie sur des méthodes bayésiennes.

Méthode de mélange Les erreurs du modèle moyen terme présenté par [21] peuvent être corrigées par ses erreurs passées pour obtenir une prévision court terme. Cette correction peut être réalisée de différentes manières, et ainsi créer de nombreux prédictors potentiels. La thèse CIFRE de Goude [54] (2008), réalisée à EDF et à l'Université d'Orsay, propose de nombreux algorithmes de mélange de prédictors et les teste en utilisant le modèle présenté par [21] pour créer des prédictors, montrant les bonnes performances des mélanges de prédictors pour cette application. Goude montre que le mélange de prédictors est efficace pour gérer les changements de périmètre du portefeuille EDF. Dans [53] (2006), Goude présente des résultats théoriques sur les mélanges de prédictors en présence de ruptures.

Faisant suite à ce travail, la thèse CIFRE de Gaillard [49], réalisée à EDF et l'Université d'Orsay et soutenue en juillet 2015, propose des algorithmes de mélange et les utilise (entre autres) sur les données du portefeuille EDF et de consommation américaine. Il utilise des experts issus de CLR, de KWF et de modèles additifs notamment. Devaine *et al.* [35] (2013) font un résumé des méthodes de mélange et les comparent sur des données de consommation française et slovaque. Gaillard *et al.* [50] (2014) proposent une méthodologie pour choisir les experts du mélange. Les méthodes de mélange étant adaptatives, elles sont efficaces pour répondre aux nouvelles problématiques d'EDF.

Le portefeuille EDF et les données françaises ont été beaucoup étudiées. Des études sur les mailles locales du réseau sont aussi réalisées au sein d'EDF.

0.1.2.2 Modélisation de la consommation locale

Méthode de sondage De Moliner [34] (2015) travaille sur 28000 courbes de charge individuelle de clients industriels, qui sont étudiées à l'aide de sondages recalés par des données n'ayant pas le même pas d'étude que les courbes de charge. Dans sa thèse CIFRE commencée en 2014 et réalisée à EDF et l'Université de Bourgogne, De Moliner étudie des méthodes d'estimation des courbes de consommation électrique moyenne pour des groupes de clients, à partir de panels de clients pour lesquels sont mesurées les courbes de charge au pas demi-horaire. Ici, l'objectif est donc d'étudier certaines courbes de charge pour en conclure des propriétés sur l'ensemble des courbes de charge.

Méthode de complétion de matrice Dans une thèse commencée en 2014 à EDF et à l'Université d'Orsay, Mei travaille sur la méthode de Nonnegative Matrix Factorization [80] pour étudier l'estimation spatio-temporelle des courbes de charge au niveau des départs HTA (niveau d'agrégation local du réseau de distribution).

Beaucoup d'études présentées précédemment portent sur la prévision court terme et ne permettent pas d'étudier l'horizon moyen terme. Les modèles additifs permettent de construire des modèles moyen et court termes efficaces en pratique.

0.1.2.3 Prévision court et moyen termes de la consommation locale et nationale

Les modèles additifs utilisant les bases de splines (voir Hastie et Tibshirani [57], 1990 et Wood [119], 2006) réalisent un bon compromis entre les modèles totalement non-paramétriques et paramétriques. Ils combinent la flexibilité des premiers et la simplicité des seconds. Des articles démontrent leur intérêt pour la prévision de la consommation électrique à un niveau national (voir e.g. Pierrot et Goude [94], 2011, Wood *et al.* [120], 2014 ou Ba *et al.* [12], 2012) et local (voir e.g. Goude *et al.* [55], 2014 ou Pompey *et al.* [96], 2015) en France, tout en demandant peu d'interventions humaines. Dans la compétition GEFCom 2012¹ présentée par Hong *et al.* [63] (2014), les modèles additifs ont été utilisés par trois des dix équipes les mieux classées (voir Nedellec *et al.* [91], 2014). La récente compétition GEFCom 2014² présentée par Hong *et al.* [64] (2015) traite de la prévision probabiliste de la charge électrique consommée, qui est un sujet en plein essor du fait du développement des énergies renouvelables et donc de l'importance des aléas, ainsi que des possibilités de gestion de la courbe de charge qu'offrent les réseaux intelligents. Les deux premières équipes Gaillard *et al.* [51] (2015) et Dordonnat *et al.* [39] (2015) ont utilisé des modèles additifs. Ces modèles sont très utilisés car ils modélisent bien la thermosensibilité de la courbe de charge et sont faciles à faire évoluer.

Dans les articles cités précédemment, le choix des covariables des modèles additifs est souvent fait en combinant expertise métier et sélection pas à pas, ce qui est chronophage et peu généralisable. La sélection et l'estimation des composantes non nulles d'un modèle additif sont donc des sujets importants. Dans [111], nous illustrons sur diverses applications de prévision de consommation l'utilisation d'une des procédures automatiques de sélection et d'estimation de composantes dans les modèles additifs établies par la suite.

0.2 Cadre statistique, estimateurs en plusieurs étapes et illustration sur des simulations

Dans le chapitre 1, nous présentons le cadre statistique dans lequel nous nous plaçons, développons les modèles additifs que nous utilisons ensuite et exprimons, ainsi que mettons à l'épreuve de simulations, plusieurs estimateurs en plusieurs étapes permettant de sélectionner et d'estimer automatiquement les composantes non nulles d'un modèle additif creux.

0.2.1 Modèle additif et méthodes de régression

Nous avons choisi d'étudier les modèles additifs [57] :

$$E(Y|X_1, \dots, X_d) = \beta_0 + \sum_{i=1}^d f_i(x_i), \quad (1)$$

où Y est la variable à expliquer (par exemple, la charge consommée), $(X_1, \dots, X_d)$ les variables explicatives (par exemple, la température ou le moment dans l'année), β_0 une constante et f_i la composante associée à la i ème variable explicative. Classiquement, nous faisons des hypothèses de régularité sur les composantes du modèle. Le modèle ainsi présenté est entièrement non-paramétrique. Son estimation est donc un problème de dimension infinie. Nous le rendons fini-dimensionnel en approchant localement les composantes dans des bases tronquées de B-Splines, réduisant ainsi le problème à un problème proche de celui des modèles linéaires en grande dimen-

1. Global Energy Forecasting Competition 2012

2. Global Energy Forecasting Competition 2014

sion. En effet, le modèle (1) peut être approché par le modèle (2) :

$$E(Y|X_1, \dots, X_d) = \beta_0 + \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(x_j), \quad (2)$$

avec $\mathbf{B}_j(-) = (B_{j,k}^{q_j}(-))_{k=1, \dots, K_j + q_j = m_j}$ la base de B-Splines de degré q_j et K_j noeuds, où est projetée la variable X_j . Soit $\mathbf{B}_j = (\mathbf{B}_j(x_{1j})^T, \dots, \mathbf{B}_j(x_{n_j})^T)^T \in \mathbb{R}^{n \times m_j}$, où x_{ij} constitue la i ème observation de la variable X_j . Les paramètres du modèle à estimer sont alors β_0 et $\boldsymbol{\beta} = (\beta_{1,1}, \dots, \beta_{d,m_d})^T$.

Nous travaillons sur la sélection de variables dans les modèles additifs. Sélectionner une covariable est équivalent à ce que sa composante soit non nulle, et donc que l'approximation de celle-ci soit non nulle, c'est-à-dire que ses coefficients estimés soient non nuls. Il est donc naturel de grouper la sélection des coefficients par covariable lors de la phase de sélection. Ceci nous conduit à utiliser le Group LASSO introduit par Yuan et Lin [122] (2006), car il permet de faire de la sélection groupe de variables par groupe de variables. Le Group LASSO a les mêmes défauts que le LASSO : en particulier, il y a un biais dans l'estimation, car il correspond à un seuillage doux des coefficients.

D'un autre côté, les estimateurs des Moindres Carrés Ordinaires (MCO) et des P-Splines (Eilers et Marx [42], 1996), qui régularisent les effets estimés en pénalisant la dérivée seconde des composantes, sont performants en estimation, mais ne permettent pas la sélection des composantes. La sous-section suivante décrit la méthodologie utilisée pour combiner le Group LASSO avec ces deux estimateurs afin de sélectionner et d'estimer convenablement les composantes additives pertinentes du modèle.

0.2.2 Estimateurs en plusieurs étapes

Nous proposons et étudions deux types d'algorithmes pour sélectionner et estimer les composantes non nulles d'un modèle additif creux.

Soit $S = \{1, \dots, d\}$ l'ensemble contenant les numéros associés aux variables explicatives présentes dans le dictionnaire de covariables. Pour une grille donnée et convenablement choisie $\Lambda_{GrpL} \in \mathbb{R}^g$, avec g dimension de Λ_{GrpL} , de valeurs de paramètre de régularisation λ associé au Group LASSO, nous utilisons, entre autres, les familles d'algorithmes Post et Ante. La seconde famille d'algorithmes sélectionne le plan d'expérience avant la phase de ré-estimation.

Nous évaluons nos procédures sur des simulations et testons Post et Ante avec trois critères de sélection de modèles (BIC, AIC, GCV). Nous travaillons sur la sélection et l'estimation de modèles additifs creux. Les plans d'expérience comportent des covariables dont les composantes sont non nulles (covariables influentes) ou nulles (covariables non influentes). Dans le cadre de simulations, nous pouvons comparer les estimations aux vrais fonctions de régression et plans d'expérience.

Pour une première étude par simulation, nous nous sommes inspirés de Lin et Zhang [83]. Il y a dix covariables, dont quatre ont des composantes non nulles dans le modèle additif. Nous pouvons modifier la corrélation entre les covariables et la variance des résidus. Ces simulations montrent la bonne capacité des procédures à identifier les covariables influentes, avec des meilleures performances pour les algorithmes de type Post, conduisant à des meilleures performances en prédiction.

Sur une seconde expérimentation, nous nous sommes inspirés du modèle présenté par [21] pour construire un simulateur de courbes de charge, qui dépend de la position dans l'année, de l'instant de la journée et d'une série de températures. Les covariables candidates pour cette simulation sont ces trois covariables, qui représentent les covariables influentes, ainsi que dix séries de températures issues de stations météorologiques différentes et une série de vitesses du vent,

Algorithme Post

1. **Première étape : Construction de sous-plans d'expérience candidats :** Pour chaque $\lambda_i \in \Lambda_{GrpL}$
 - Notons $S_{\lambda_i} \subset S$ le sous-ensemble des numéros des covariables sélectionnées par l'estimateur associé au Group LASSO de paramètre λ_i
 - Résoudre $\tilde{\beta}_{\lambda_i} = \arg \min_{\beta} \sum_{l=1}^n \left(Y_l - \beta_0 - \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(X_{l,j}) \right)^2 + \lambda_i \sum_{j=1}^d \sqrt{m_j} \|\beta_j\|_2$
 - Pour chaque $j \in S$, si $\|\tilde{\beta}_{\lambda_i,j}\|_2 \neq 0$, ajouter j dans S_{λ_i} , sinon ne pas considérer j dans S_{λ_i} . Pour faciliter les notations, nous supposons que tous les sous-plans d'expérience candidats sont différents.
 2. **Seconde étape : Estimation des modèles candidats :** Pour chaque $S_{\lambda_s} \in \{S_{\lambda_{min}}, \dots, S_{\lambda_{max}}\}$
 - Résoudre un problème du type $\hat{\beta}_{S_{\lambda_s}} = \arg \min_{\beta} Q(\beta)$ où Q est la fonction objective des MCO ou des P-Splines
 - Calcul du critère de sélection de modèle (CSM), typiquement le BIC [102], l'AIC [1] ou le GCV [115], de l'estimateur $\hat{\beta}_{S_{\lambda_s}}$
 3. **Troisième étape : Sélection de l'estimateur final :** Sélectionner $\hat{\beta}_{S_{\lambda_b}}$ qui minimise le CSM choisi.
-

Algorithme Ante

1. **Première étape : Construction de sous-plans d'expérience candidats :** Pour chaque $\lambda_i \in \Lambda_{GrpL}$
 - Notons $S_{\lambda_i} \subset S$ le sous-ensemble des numéros des covariables sélectionnées par l'estimateur associé au Group LASSO de paramètre λ_i
 - Résoudre $\tilde{\beta}_{\lambda_i} = \arg \min_{\beta} \sum_{l=1}^n \left(Y_l - \beta_0 - \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(X_{l,j}) \right)^2 + \lambda_i \sum_{j=1}^d \sqrt{m_j} \|\beta_j\|_2$
 - Pour chaque $j \in S$, si $\|\tilde{\beta}_{\lambda_i,j}\|_2 \neq 0$, ajouter j dans S_{λ_i} , sinon ne pas considérer j dans S_{λ_i} . Pour faciliter les notations, nous supposons que tous les sous-plans d'expérience candidats sont différents.
 2. **Seconde étape : Sélection du plan d'expérience final :** Sélectionner S_{λ_b} tel que $\tilde{\beta}_{\lambda_b}$ minimise le CSM choisi.
 3. **Troisième étape : Estimation de l'estimateur :** Résoudre un problème du type $\hat{\beta}_{S_{\lambda_b}} = \arg \min Q$ où Q est la fonction objective des MCO ou des P-Splines.
-

qui constituent les variables non influentes. Comme pour les premières simulations, les méthodes de type Post sont plus efficaces en termes de sélection et de prévision.

0.2.3 Perspectives méthodologiques

Comme le LASSO, le Group LASSO a tendance à ne sélectionner qu'un seul groupe de covariables lorsque ceux-ci sont fortement corrélés. Une pénalisation tenant compte de la corrélation peut être utilisée à place du Group LASSO. Par exemple, les auteurs de [43] combinent une pénalité L1 avec une pénalité tenant compte de la corrélation dans le contexte linéaire.

L'ajout de nouvelles données pose la question de l'adaptativité des méthodes. Avec une bonne initialisation, les procédures peuvent être en partie adaptatives.

Nous avons travaillé sur les modèles additifs. Un travail similaire sur les modèles à coefficients variables peut être réalisé.

0.3 Estimateurs en plusieurs étapes de modèles additifs appliqués à la prévision de consommation électrique à plusieurs niveaux d'agrégation

Nous travaillons dans le chapitre 2 sur trois jeux de données de prévision de consommation électrique : le portefeuille EDF, comportant les consommations des clients d'EDF et étant donc à un haut niveau d'agrégation, les données GEFCOM 2012 [63] disponibles via le site Kaggle (données locales américaines issues d'une compétition publique) et des données de consommation d'un niveau local d'une maille du réseau électrique français, les postes sources.

Nous présentons d'abord l'intérêt de ces trois jeux de données.

Le portefeuille EDF est directement influencé par les changements des habitudes des consommateurs et par l'ouverture du marché électrique. Du fait de ces évolutions, EDF a besoin de modèles et de méthodes qui s'adaptent rapidement, notamment pour modéliser la thermosensibilité de la courbe de charge. Le modèle additif répond efficacement à ce besoin. Nous vérifions empiriquement que nos procédures sélectionnent efficacement les covariables météorologiques. De plus, comme nous connaissons bien ces données, nous pouvons facilement tester de nouvelles méthodes, avant de les appliquer sur d'autres données moins connues.

Les applications sur les deux autres jeux de données ont des objectifs communs. Nous cherchons à sélectionner automatiquement les localisations géographiques où sont mesurées les covariables météorologiques d'un modèle additif pour avoir les erreurs de prévision les plus faibles possibles. Étudier les postes sources est très intéressant car ils constituent un sujet actuel et prenant pour ERDF. Cependant, ceux-ci présentent le défaut majeur d'être hautement confidentiels. Nous ne pouvons donc pas entièrement les exploiter ni les présenter. Ceci nous a conduit à étudier les données de GEFCOM 2012, qui correspondent en partie à une idéalisation du problème posé par les postes sources. Ce n'est par ailleurs pas le seul intérêt de ces données. Ce type de compétitions est un formidable outil pour créer une émulation scientifique importante autour de quelques problématiques portant sur la prévision de consommation électrique. Elles offrent un socle commun de travail aux prévisionnistes du monde entier, ce qui est rare dans un domaine où les données sont souvent confidentielles. Elles conduisent à une meilleure communication des méthodes et à la possibilité de nous comparer à des procédures externes à EDF.

0.3.1 Application agrégée au niveau national

Cette application est explicitée dans [111]. Le portefeuille EDF comprend plusieurs segments des clients d'EDF : les profilés (petits clients, entreprises moyennes ou grandes, dont le type de comptage est profilé), les 32000 (grandes entreprises avec un comptage télérelevé), les ELD (Entreprise Locale de Distribution d'énergie électrique, télérelevé) et les sup7 (entreprises spéciales, télérelevé). Ce signal ne comprend pas les pertes réseaux (pertes RTE ou ERDF), les auto-consommations, les clients perdus, les effacements tarifaires et contractuels ainsi que les échanges. Il s'agit d'un niveau hautement agrégé de la consommation électrique française. À l'opérationnel, ce signal est utilisé dans des départements tels que la Direction Commerce opérationnelle, particulièrement touchée par la fin à venir des tarifs réglementés, ou la Direction Optimisation Amont Aval et Trading (DOAAT), dont la mission principale est de maximiser la marge brute d'électricité d'EDF, à tout horizon de temps, en respectant les risques fixés par la Direction générale et en laissant à l'amont et l'aval ses leviers propres.

Les modèles additifs s'adaptent efficacement aux variations de ce portefeuille. La modélisation additive a été améliorée sur des données similaires. La DOAAT a demandé à un expert d'EDF R&D, avec qui nous avons travaillé, d'optimiser un modèle additif de prévision du portefeuille EDF. L'estimation de ce nouveau modèle s'intègre dans un projet de développement de ces modèles à EDF.

L'expert d'EDF R&D a construit un dictionnaire de covariables calendaires et météorologiques (instantanées, agrégées ou lissées) de grande taille. Pour sélectionner son modèle additif final, l'expert teste un nombre important de modèles plausibles. La méthode est longue, fastidieuse et ne peut être facilement généralisée à d'autres jeux de données. Ceci motive l'application de procédures automatiques. Nous nous focalisons sur la sélection de variables météorologiques. Nous travaillons à la fois sur du court et du moyen termes. L'objectif de l'étude n'est pas d'obtenir des modèles ayant des performances prédictives bien supérieures à celles du modèle de l'expert, car celui-ci est déjà très performant, mais consiste à obtenir des performances au moins équivalentes sans intervention humaine, tout en vérifiant la cohérence physique des variables sélectionnées et en cherchant les origines des erreurs. Nous montrons que cet objectif est atteint par nos méthodes.

0.3.2 Application locale

0.3.2.1 GEFCOM 2012

Cette application est aussi explicitée dans [111]. GEFCOM 2012 [63] est une compétition mise en ligne par le site Kaggle qui propose de prévoir et de reconstruire la charge horaire consommée (exprimée en kW) de 20 zones américaines. En plus des courbes de charge, les concurrents ont à disposition les températures de onze stations météorologiques entre 2004 et juin 2008. La structure du réseau est cachée : les participants n'ont ni les localisations des zones de consommation électrique à prévoir, ni celles des stations météorologiques. L'objectif initial de la compétition est de prévoir la dernière semaine de la base de données et de reconstruire sept semaines en 2005. Nous utilisons ces données pour prévoir les six premiers mois de 2008.

La Figure 3 donne une représentation factice de l'aspect géographique de la compétition. La station météorologique la plus proche d'une station électrique n'est pas nécessairement la plus proche de tous les groupes de consommateurs reliés à la station électrique. Chaque zone peut ainsi potentiellement couvrir un grand territoire à cause de la topologie du réseau (qui est inconnue). Les participants de la compétition doivent construire une modélisation malgré la structure cachée du réseau. Sur des applications propres à EDF, la structure du réseau peut aussi être cachée pour certaines études pour des raisons de confidentialité.

Ici, plusieurs notions propres à la prévision de la consommation électrique sont abordées. Dans GEFCOM 2012, plusieurs horizons de prévision (du journalier à l'hebdomadaire) ainsi que

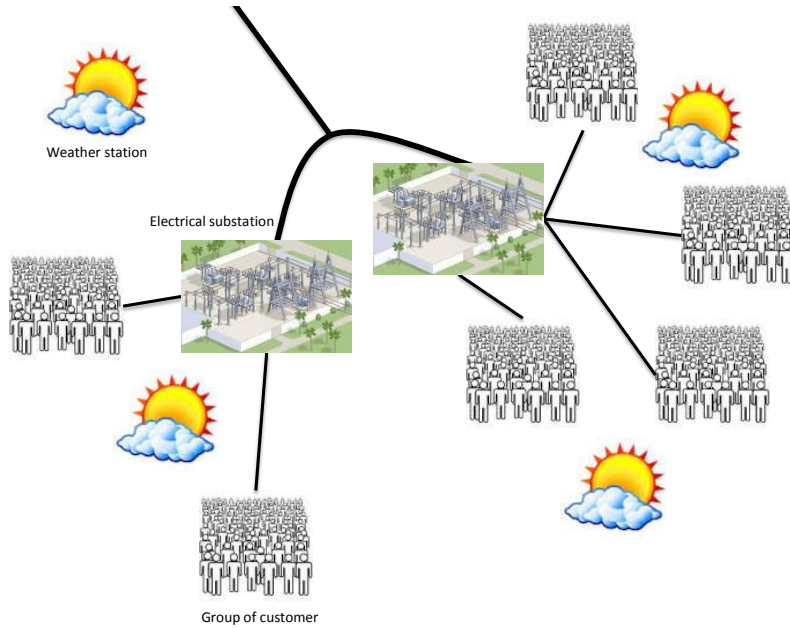


FIGURE 3 – Illustration de la compétition GEFCom 2012

la reconstruction de la courbe de charge, qui est un sujet important pour EDF concernant les effacements par exemple, sont étudiés. L'intégration et la simulation de la température sont aussi traitées. Il s'agit d'une étude locale, ce qui implique la présence de signaux de consommation plus variables et conduit à un besoin d'automatisation des méthodes. L'absence d'information géographique impose de travailler sur le sujet, souvent peu abordé, de la sélection de stations météorologiques et permet ainsi d'en étudier l'impact.

Un problème important de l'application est de sélectionner les points où sont mesurées les températures utilisées dans les modèles. Bien que les stratégies aient été différentes, les quatre premières équipes (voir [84], [105], [91] et [17]) de la compétition GEFCom 2012 ont proposé un nombre fixe de stations météorologiques (1, 5 ou 11). Après la compétition, Hong *et al.* [65] montrent que la prévision peut être améliorée en relâchant cette hypothèse. Nos procédures automatiques sont une solution pour répondre à ce problème. Nous montrons empiriquement que la sélection des stations météorologiques est importante malgré la forte corrélation entre les températures et l'introduction de plusieurs signaux de températures, et non un unique, dans les modèles de consommation électrique permet d'améliorer la prévision. Enfin, nous montrons que nos procédures sont performantes pour cette application.

0.3.2.2 Postes sources

Avec la création du marché européen de l'électricité, EDF a été contraint de séparer ses activités de production, de transport et de distribution d'électricité. Dans ce cadre, RTE a vu le jour en 2000. Cette filiale d'EDF a pour mission d'assurer un accès équitable au réseau de transport à tous les acteurs du marché de l'électricité en France. ERDF, autre filiale d'EDF, créée en 2008, gère l'acheminement et la distribution de l'électricité aux clients d'EDF. La liaison entre les réseaux de transport et de distribution a lieu au niveau des postes sources. Au nombre d'environ 2200, ils sont gérés par ERDF. Alimenté par le réseau de transport haute tension, le réseau de distribution a un mode de fonctionnement radial : l'électricité circule des postes sources (postes HT/MT) en amont, qui abaissent la tension par une succession de lignes et de transformateurs, vers les installations des consommateurs en aval. Il existe localement des petites sources de production (éoliennes, micro-centrales hydrauliques, photovoltaïques, etc) qui injectent

de l'électricité sur le réseau.

Pour gérer la grille de distribution, quantifier les contraintes et optimiser la configuration du réseau, ERDF a besoin de procédures permettant de prévoir la charge consommée à court et moyen termes pour chaque poste source et est dans l'obligation de fournir chaque année une prévision de la charge consommée pendant l'année à venir pour tous les postes sources. En cas d'erreurs importantes dans les prévisions, en plus des risques de panne sur le réseau, ERDF s'expose à des pénalités financières. Plus de 2200 postes sources, et donc plus de 2200 séries chronologiques, sont à étudier. En moyenne, chaque poste source est relié à 40 gros et 16000 petits consommateurs. Cependant, les profils des postes sources peuvent complètement changer selon les types de consommateurs qui lui sont reliés (un poste source relié principalement à des gros clients industriels est plus dépendant du moment dans l'année par exemple) ainsi que de la localisation (un poste source dans le Sud risque d'être plus sensible aux fortes chaleurs car il y a plus de climatiseurs installés par exemple). Nous avons travaillé sur une soixantaine de postes sources de l'ACR³ Lyon. Comme pour GEFCOM 2012, nous nous sommes focalisés sur la sélection des stations météorologiques. Celle-ci améliore les performances de la prévision de la consommation électrique pour les postes sources de cette ACR.

0.3.3 Résumé des résultats et implémentation des méthodes

Résumé des résultats Nous avons travaillé sur la modélisation moyen et court termes de la consommation électrique sur plusieurs jeux de données, illustrant les bonnes propriétés empiriques des modèles additifs pour modéliser la consommation électrique à différents niveaux d'agrégation et d'horizon.

Nous nous sommes plus particulièrement intéressés à la sélection de variables dans les modèles additifs pour en améliorer la prévision. Jusqu'à présent, les méthodes utilisées à EDF pour sélectionner les covariables d'un modèle additif de prévision de consommation électrique sont principalement manuelles, ce qui est peu adapté aux changements actuels du paysage électrique français et au besoin d'automatisation induit par le développement des études des mailles locales. Notre étude montre que, sur trois jeux de données différents de prévision de consommation électrique, une de nos procédures permet de sélectionner automatiquement les variables d'un modèle additif en ne demandant ni expertise métier, ni intervention humaine, avec des bonnes performances prédictives. Nous avons plus particulièrement étudié les bonnes capacités à modéliser la thermosensibilité par les modèles additifs.

Nous avons travaillé sur la correction court terme de modèles moyen terme en utilisant des auto-regressifs comportant ou non une moyenne mobile. Bien que cela n'apparaisse pas dans la suite du document, que nous ne souhaitons pas alourdir, nous avons constaté que notre correction obtient de meilleures performances qu'un correctif à poids constant utilisé à EDF ou qu'un correctif à l'aide de forêts aléatoires.

Nous avons travaillé sur le sujet portant sur la sélection des localisations où sont mesurées les covariables météorologiques introduites dans le modèle additif pour l'étude d'une maille locale. La stratégie consistant à sélectionner la station la plus proche du point étudié de la maille du réseau électrique n'est pas nécessairement la plus efficace. Ces études nous ont aussi permis d'illustrer le bon comportement des modèles additifs malgré une corrélation forte entre les covariables. Nous montrons les bonnes capacités d'une de nos procédures à automatiser la sélection et l'estimation de composantes non nulles d'un modèle additif sur un nombre important de séries chronologiques, sans nécessiter d'intervention humaine et en un temps raisonnable.

Enfin, nos procédures fabriquent un nombre important de prédicteurs, qui pourraient ensuite être candidats dans des modèles de mélange.

Prototype de package R CASA Le prototype de package R **CASA**, présenté à la conférence UseR! 2015, implémente la méthodologie utilisée pour ces applications. Le prototype de package, qui est présenté dans une vignette dans l'Annexe A, comprend des fonctions réalisant la sélection et l'estimation automatique de composantes non nulles de modèles additifs. Il permet d'étudier des séries temporelles simples et multiples, permettant ainsi de traiter certaines particularités que possède la consommation électrique, mais qui se retrouvent sur d'autres applications. Dans la vignette, nous proposons un exemple d'application du code portant sur les données de GEFCom 2012.

0.3.4 Autres applications

Mentionnons trois exemples d'application pour lesquelles nous aurions pu utiliser nos procédures.

Dans l'article [70], Jollois *et al.* (2009) proposent d'utiliser les modèles additifs pour prévoir la pollution en Haute-Normandie. Pour cela, [70] introduit un certain nombre de covariables météorologiques dépendant de la température, de l'humidité, de la pluie et du vent. Comme pour l'étude du portefeuille EDF, nous aurions pu créer un dictionnaire de covariables météorologiques issues de ces mesures pour ensuite en sélectionner certaines.

Dans une étude interne à EDF, des modèles additifs ont été appliqués à la prévision de production d'un champ d'éoliennes. La problématique posée par l'étude est proche de celle de GEFCom 2012. L'objectif est de trouver les signaux de vitesses du vent permettant la meilleure prévision de la production. Pour cela, le vent est mesuré le long d'une grille. La station météorologique la plus proche du champ conduit à une moins bonne prévision que l'agrégation de l'ensemble des vents de la grille, ce que nous expliquons par la variabilité du vent.

Marra et Wood [86] (2011) utilisent les modèles additifs pour étudier la relation entre la concentration de bêta-carotènes dans le plasma et des caractéristiques personnelles. Le jeu de données à disposition contient beaucoup de covariables continues.

0.4 Un résultat de consistance pour un estimateur en plusieurs étapes

Nous présentons le chapitre 3 de la thèse, qui est consacré à l'étude asymptotique de l'une des procédures Post.

Le problème de sélection de composantes est un sujet classique en Statistiques. Dans le cas des modèles additifs, Fan et Jiang [44] utilisent des tests de déviance. Nous montrons sur une simulation dans le chapitre 1 qu'une corrélation élevée entre les covariables peut fortement réduire leur efficacité. Les méthodes pénalisées, que nous utilisons, sont connues pour avoir un coût informatique plus faible et pour être plus robustes aux fortes corrélations présentes dans le plan d'expérience (voir [112] ou [7]). Marra et Wood [86] proposent une version modifiée des P-Splines [42]. Nous constatons sur des simulations que la norme 2 des coefficients estimés des effets non influents n'est pas strictement nulle. Meier *et al.* [89] proposent d'utiliser la pénalité (3) :

$$\lambda_1 \sum_{j=1}^d \sqrt{\|f_j\|_{2_n}^2 + \lambda_2 I(f_j)^2}, \quad (3)$$

où $\|f_j\|_{2_n}^2 = 1/n \sum_{i=1}^n f_j(X_{i,j})^2$, permettant de sélectionner les composantes non nulles, et $I(f_j)^2 = \int_0^1 f_j^{(2)}(x)^2 dx$, permettant de régulariser les estimateurs. Nous constatons sur une simulation inspirée de [83] que la présence de covariables non influentes dans le plan d'expérience dégrade la performance en prédiction de l'estimateur comparée à celle du même estimateur lorsque le plan

d'expérience ne contient que les covariables influentes, et ce même si l'estimateur permet bien de ne sélectionner que les covariables influentes. Ceci justifie la séparation de la phase de création des sous-plans d'expérience candidats et la phase d'estimation. Des auteurs utilisent des estimateurs en plusieurs étapes pour avoir des méthodes efficaces à la fois en termes de sélection et d'estimation. Antoniadis *et al.* [9] utilisent un estimateur Nonnegative Garrote [20] ayant pour estimateur initial un estimateur P-Splines. Huang *et al.* [67] utilisent un estimateur Group LASSO adaptatif, dont l'estimateur initial est un estimateur Group LASSO [122].

Nous supposons que les fonctions f_j du modèle (1) appartiennent à un espace de Sobolev. Pour approcher le modèle, nous projetons les composantes additives dans des bases tronquées de B-Splines [33]. Nous nous ramenons donc à un contexte proche du contexte linéaire. Sélectionner une covariable revient à sélectionner le groupe de coefficients qui lui est associé. Nous avons choisi d'utiliser le Group LASSO. Celui-ci permet de réaliser la sélection des composantes, mais a les mêmes défauts que le LASSO [112]. A cause de la pénalisation L1, les coefficients estimés, même les plus forts, sont sur-rétrécis, conduisant à une estimation des composantes fortement biaisée. Le Group LASSO correspond à un seuillage doux. Le degré de pénalisation constant quelle que soit l'amplitude des coefficients explique en partie le biais. Pour compenser cet effet, le Group LASSO a tendance à inclure des variables non pertinentes. Pour corriger ce biais, et pour mettre au point des méthodes simultanément optimales en sélection et en prédiction asymptotique, nous pouvons utiliser des estimateurs en plusieurs étapes. Nous avons trouvé peu de résultats théoriques démontrés concernant les estimateurs en plusieurs étapes. Dans le contexte linéaire, Belloni *et al.* [16] ont considéré une méthode de type Ante, remplaçant le Group LASSO par du LASSO et en utilisant des MCO pour la phase d'estimation. Ils prouvent que l'estimateur obtenu a des performances au moins aussi bonnes en termes de vitesse de convergence que le LASSO et permet de réduire le biais. Antoniadis *et al.* [8] appliquent des variantes de nos procédures avec des bons résultats pratiques. Ceci motive notre étude. Dans le chapitre 3, nous étudions plus en détails l'estimateur de type Post1 qui combine le Group LASSO avec l'estimateur MCO.

Nous nous plaçons dans un contexte asymptotique pour lequel le nombre d'observations tend vers l'infini et le nombre de covariables candidates peut éventuellement être plus élevé que le nombre d'observations disponibles, mais avec moins de covariables "influentes" que d'observations. Nous nous focalisons sur le Group LASSO et les B-Splines. Plus précisément, nous devons tenir compte du fait que :

- Chaque composante additive f_j du modèle est approchée par une combinaison linéaire de fonctions de bases (les B-Splines), leur nombre pouvant croître avec le nombre d'observations n pour que l'approximation soit de bonne qualité.
- Nous appliquons le Group LASSO aux vecteurs des coefficients résultants de cette approximation puis nous ré-appliquons du Group LASSO ou des MCO ou des P-Splines pour améliorer la sélection (dans les cas des procédures de type Post) et l'estimation en dé-biaisant l'estimation par Group LASSO.

Pour notre travail, la notion d'effet significatif d'une variable ne se traduit pas, comme il est usuel dans ce contexte, par des normes des composantes non nulles du modèle bornée inférieurement par une constante strictement positive. Nous supposons que la norme de chaque composante significative est minorée par une suite décroissante dépendant du nombre d'observations et pouvant tendre vers 0 asymptotiquement. Nous nous sommes inspirés du travail de Wang *et al.* [116] qui montrent la consistance en sélection du BIC dans le contexte linéaire. Après avoir montré la consistance d'un critère type BIC dans le contexte additif lorsque les composantes sont approximées dans des bases de B-Splines, nous montrons la consistance en sélection de ce que nous appelons Post1Bic puis utilisons un résultat de Kato [72] pour montrer la consistance en estimation. Le chapitre 3 a été publié dans [10] et une synthèse a été soumise dans les actes de conférence de l'ISNPS 2014 dans [6].

Nous avons montré des propriétés de consistance pour Post1Bic. L'étude de la procédure de

type Post combinant le Group LASSO avec les P-Splines est plus compliquée, car l'estimateur dépend de la pénalisation associée à l'estimateur P-Splines. Notre méthode de démonstration ne permet pas de conclure des propriétés de consistance lorsque l'AIC est utilisé sous nos hypothèses, car sa pénalisation n'est pas suffisamment forte, conduisant à des faux négatifs.

Cadre statistique, estimateurs en plusieurs étapes et illustration sur des simulations

Introduction

L'électricité se stocke mal (voir [87]). EDF doit donc toujours être à l'équilibre entre production et consommation d'électricité. Ceci crée un besoin de modélisation et de prévision de consommation d'électricité à différents horizons, tant à court terme (jour pour le lendemain), à moyen terme (prévision à un horizon de quelques semaines à une année) qu'à long terme (prévision de plusieurs années). Plusieurs évolutions majeures, présentées en introduction, sont venues s'ajouter à cette nécessité d'équilibre. Ces évolutions conduisent à l'étude de nombreuses séries chronologiques (plus de 2200 postes sources), pouvant avoir des comportements différents, et au besoin de nouvelles méthodes s'adaptant rapidement. Des méthodes de sélection de variables automatiques sont donc nécessaires.

Les problèmes industriels que rencontre EDF sont complexes et demandent des réponses statistiques adéquates. Nous introduisons dans ce chapitre les modèles, les méthodes et les algorithmes utilisés pour construire un cadre statistique répondant à ces problèmes. Pour plus de simplicité, nous travaillons d'abord sur les modèles linéaires, puis analysons les modèles additifs. Plus précisément, dans la section 1.1, nous introduisons des notions fondamentales sur les modèles linéaires, ainsi que ses estimateurs les plus classiques (estimateur par Moindres Carrés Ordinaires, Ridge, LASSO...). Nous justifions l'intérêt de la sélection de composantes en expliquant le principe de parcimonie. Nous présentons des méthodes classiques de sélection de variables et justifions notre choix d'utiliser des méthodes de régression pénalisée. Dans la section 1.2, nous présentons les modèles additifs, proposés par Hastie et Tibshirani [57], ainsi que les bases de B-Splines (De Boor, [33]), utilisées pour approximer les composantes du modèle additif. Nous introduisons ensuite une méthode de placement des noeuds des bases de B-Splines due à de Zhou et Shen [123]. Dans la section 1.3, nous introduisons les méthodes de régression pénalisée associées à l'estimation de modèles additifs, et notamment le Group LASSO (Yuan et Lin, [122]) et les P-Splines (Eilers et Marx, [42]), puis nous expliquons les méthodes que nous utilisons pour les combiner. Dans la section 1.4, nous menons une étude par simulation inspirée de Lin et Zhang [83] et une pseudo-simulation, inspirée d'un modèle de prévision de consommation classiquement utilisé à EDF [21].

1.1 Généralités sur le modèle linéaire

1.1.1 Définition

Nous cherchons à expliquer et à prédire une variable aléatoire Y à l'aide des variables aléatoires $\mathbf{X} = (X_1, \dots, X_d)$, où $d > 0$ est le nombre de covariables (variables explicatives). Nous disposons de

$\mathbf{y} = (y_i)_{i \in 1, \dots, n}$ et $\mathbf{x} = (\mathbf{x}_i)_{i \in 1, \dots, n}$ réalisations d'un n -échantillon du vecteur aléatoire $(Y_i, \mathbf{X}_i)$, chaque $d + 1$ uplet étant une copie indépendante du $d + 1$ uplet de $(Y, \mathbf{X})$. Posons $\mathbf{Y} = (Y_1, \dots, Y_n)$ et $\mathbf{X} = (\mathbf{X}_1^T, \dots, \mathbf{X}_n^T)$. Nous cherchons une fonction f mesurable telle que $f(\mathbf{X})$ "approche au mieux" Y . Si les moments d'ordre 2 de $(Y, \mathbf{X})$ existent, "approcher au mieux" peut correspondre à trouver f mesurable qui minimise $E((Y - f(\mathbf{X}))^2)$, minimisant ainsi l'erreur quadratique moyenne. La solution est alors $E(\mathbf{Y}|\mathbf{X})$. La première simplification possible consiste à considérer cette espérance conditionnelle comme linéaire en $\mathbf{X}$, c'est-à-dire que la forme de f est linéaire. $\mathbf{Y}$ suit un modèle linéaire si

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}, \quad (1.1)$$

où $\boldsymbol{\beta} \in \mathbb{R}^d$ inconnu (à estimer), $\boldsymbol{\varepsilon} = (\varepsilon_1, \dots, \varepsilon_n)$ vecteur aléatoire tel que $(\varepsilon_i)_i$ i.i.d., centrés et de variance σ^2 . Le modèle linéaire est gaussien si $\mathbf{Y}$ est gaussien. Nous noterons $\mathbf{Y} \sim N(\mathbf{X}\boldsymbol{\beta}, \sigma^2 I_n)$.

Il se pose alors les trois problèmes statistiques suivants :

- Problème d'estimation : estimer $\boldsymbol{\beta}$ et σ^2
- Problème de prédiction : estimer $\mathbf{X}\boldsymbol{\beta}$, avoir une nouvelle valeur $\hat{Y}$ à partir de $\mathbf{X}^{(1)}\boldsymbol{\beta}$ avec $\mathbf{X}^{(1)}$ nouvelle copie de $\mathbf{X}$
- Problème de sélection : estimer $\{j | \beta_j \neq 0\}$

1.1.2 Un estimateur classique : l'estimateur des moindres carrés ordinaires

"Approcher au mieux" dans un modèle linéaire correspond pour nous à minimiser $E((Y - \mathbf{X}\boldsymbol{\beta})^2)$, qui n'est calculable que si la loi $(\mathbf{Y}, \mathbf{X})$ est connue, ce qui n'est généralement pas le cas. Pour les modèles linéaires, une version empirique du critère est $\|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|_2^2$. Il est alors usuel de considérer l'estimateur des moindres carrés ordinaires (EMCO) pour minimiser ce critère empirique. L'EMCO de $\boldsymbol{\beta}$ pour le modèle (1.1) est défini par l'équation (1.2) :

$$\hat{\boldsymbol{\beta}}^{MCO} \in \arg \min_{\boldsymbol{\beta} \in \mathbb{R}^d} \|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|_2^2, \quad (1.2)$$

qui minimise la fonction objective

$$Q^{MCO}(\boldsymbol{\beta}) = \|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|_2^2.$$

$\|\cdot\|_2$ correspond à la norme euclidienne. L'EMCO est la projection orthogonale de $\mathbf{Y}$ sur l'espace vectoriel engendré par les colonnes de $\mathbf{X}$.

Nous pouvons écrire le système des équations normales :

$$(\mathbf{X}^T \mathbf{X}) \hat{\boldsymbol{\beta}}^{MCO} = \mathbf{X}^T \mathbf{Y}. \quad (1.3)$$

Pour identifier une solution unique, $\mathbf{X}^T \mathbf{X}$ doit être de plein rang, et donc $\text{rang}(\mathbf{X}) = d$, c'est-à-dire que $\mathbf{X}$ doit être injective. La solution de (1.3) est alors donnée par l'expression suivante :

$$\hat{\boldsymbol{\beta}}^{MCO} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}.$$

Posons $H = \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$, alors $\hat{\mathbf{Y}} = \mathbf{X} \hat{\boldsymbol{\beta}}^{MCO} = H\mathbf{Y}$. H est la matrice de projection sur le sous-espace engendré par les colonnes de $\mathbf{X}$. Soit $\mathbf{Z} = (I - H)\mathbf{Y}$. Alors $\mathbf{Z} \sim N(0, \sigma^2(I - H))$ et $E(\|\mathbf{Z}\|^2) = \sigma^2 \text{tr}(I - H)$, où tr désigne l'opérateur trace. Le nombre de degrés de liberté est $n - \text{tr}(H)$.

D'après le théorème de Gauss-Markov, l'EMCO est l'unique estimateur linéaire sans biais de variance minimale parmi les estimateurs linéaires sans biais de $\boldsymbol{\beta}$. La variance des estimateurs est égale à :

$$\text{var}(\hat{\boldsymbol{\beta}}^{MCO}) = \sigma^2 (\mathbf{X}^T \mathbf{X})^{-1}.$$

$\underline{\mathbf{X}}^T \underline{\mathbf{X}}$ est une matrice symétrique réelle définie positive. Il existe P matrice orthogonale et $(\mu_1, \dots, \mu_d)$ réels tels que

$$\underline{\mathbf{X}}^T \underline{\mathbf{X}} = P^T \text{diag}(\mu_1, \dots, \mu_d) P,$$

où $\text{diag}(\mu_1, \dots, \mu_d)$ est la matrice diagonale de termes diagonaux $(\mu_1, \dots, \mu_d)$, qui sont les valeurs propres de $\underline{\mathbf{X}}^T \underline{\mathbf{X}}$. Alors

$$(\underline{\mathbf{X}}^T \underline{\mathbf{X}})^{-1} = P \text{diag}(1/\mu_1, \dots, 1/\mu_d) P^T.$$

La variance des estimateurs de certaines composantes est élevée lorsque les valeurs propres correspondantes de $\underline{\mathbf{X}}^T \underline{\mathbf{X}}$ sont petites. Comment remédier à ce problème ?

1.1.3 Estimateur Ridge

Un mauvais conditionnement de $\underline{\mathbf{X}}$ conduit à la présence de valeurs propres faibles et donc à une variance de l'EMCO élevée. Une petite variation de $\mathbf{Y}$ peut alors conduire à une forte variation de l'EMCO. Ce manque de robustesse peut être limité en imposant des contraintes sur les coefficients dans le critère d'optimisation conduisant à l'EMCO. Nous obtenons ainsi l'estimateur Ridge [62] :

$$\hat{\boldsymbol{\beta}}^{\text{Ridge}} \in \arg \min_{\boldsymbol{\beta}} Q^{\text{MCO}}(\boldsymbol{\beta}) \text{ s.c. } \|\boldsymbol{\beta}\|_2^2 \leq t. \quad (1.4)$$

Pour le critère (1.4), la norme euclidienne des coefficients est inférieure à un certain seuil $\sqrt{t}$, limitant ainsi l'impact d'une variation de $\mathbf{Y}$. L'estimateur est plus robuste.

La formule lagrangienne de l'équation (1.4) admet une autre interprétation : il s'agit d'une version pénalisée du problème des MCO.

$$\hat{\boldsymbol{\beta}}^{\text{Ridge}} \in \arg \min_{\boldsymbol{\beta}} Q^{\text{MCO}}(\boldsymbol{\beta}) + \lambda \sum_{i=1}^d \beta_i^2, \quad (1.5)$$

avec $\lambda \geq 0$. L'estimateur Ridge s'écrit alors :

$$\hat{\boldsymbol{\beta}}^{\text{Ridge}} = (\underline{\mathbf{X}}^T \underline{\mathbf{X}} + \lambda \mathbf{I}_d)^{-1} \underline{\mathbf{X}}^T \mathbf{Y}. \quad (1.6)$$

Les valeurs propres de $\underline{\mathbf{X}}^T \underline{\mathbf{X}} + \lambda \mathbf{I}_d$ sont $(\mu_1 + \lambda, \dots, \mu_d + \lambda)$. Elles sont strictement positives. Un choix approprié de λ modifie le conditionnement de $\underline{\mathbf{X}}^T \underline{\mathbf{X}}$. Le risque quadratique de $\hat{\boldsymbol{\beta}}^{\text{Ridge}}$ est donné par l'expression suivante :

$$\sigma^2 \sum \mu_i / (\mu_i + \lambda)^2 + \lambda \boldsymbol{\beta}^T (\underline{\mathbf{X}}^T \underline{\mathbf{X}} + \lambda \mathbf{I}_d)^{-1} \boldsymbol{\beta}.$$

Si $\lambda = 0$, le risque quadratique est $\sigma^2 \sum 1/\mu_i$. L'estimateur est alors le même que l'EMCO : il est de biais nul, mais de variance plus élevée que lorsque $\lambda \neq 0$. Lorsque λ croît, le biais augmente mais la variance diminue. Enfin, lorsque λ tend vers l'infini, l'estimateur Ridge tend vers 0 : la variance et le biais sont respectivement nulle et élevé. Le paramètre λ contrôle donc un compromis entre le biais et la variance.

L'estimateur Ridge répond au problème statistique d'estimation. Un des problèmes statistiques classiques associés à la régression est d'estimer $\{j | \beta_j \neq 0\}$, c'est-à-dire de réaliser de la sélection de covariables.

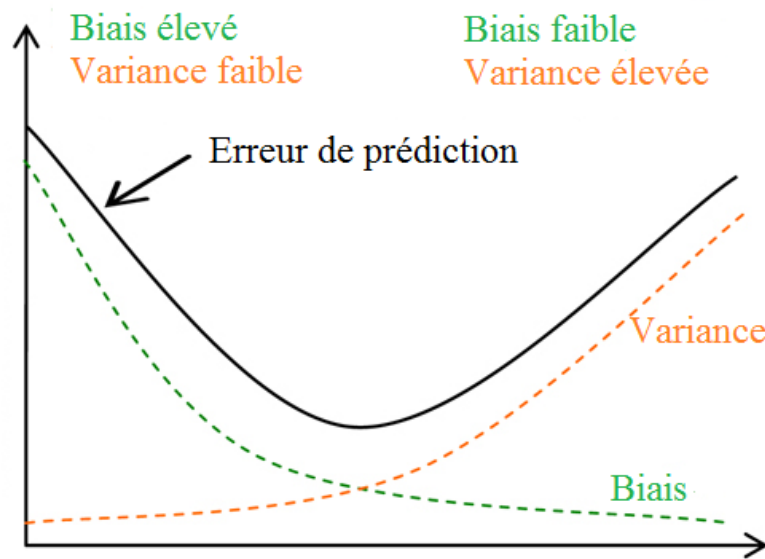


FIGURE 1.1 – Principe de parcimonie

1.1.4 Généralités sur la sélection de variables

Lorsqu'un modèle est fixé, la théorie offre un cadre rigoureux s'assurant la construction d'estimateurs performants. Cependant, dans la majorité des applications pratiques, le modèle est inconnu, ce qui conduit à la problématique courante en Statistiques de sélection de modèles. La sélection de variables est un cadre classique de sélection de modèles. Nous cherchons alors à sélectionner les variables explicatives qui ont le plus d'effet sur la variable à expliquer, et ainsi à améliorer l'interprétabilité et l'adéquation aux données tout en s'assurant de la parcimonie. La réalisation du principe de parcimonie (voir Figure 1.1) peut être vue comme l'optimisation en termes d'erreur de prédiction du compromis biais/variance. En général, le biais diminue et la variance augmente lorsque le nombre de paramètres augmente. En effet, un modèle sous-paramétré ne prend pas en compte des effets importants. Le biais des estimateurs est alors élevé et la variance sous-estimée. Inversement, lorsque le modèle est sur-paramétré, le biais est généralement nul, mais la variance des estimateurs est plus élevée et il y a un risque de sur-apprentissage, qui est néfaste pour la prédiction.

La sélection de variables peut répondre à plusieurs objectifs, que nous énumérons ensuite :

- améliorer la prédiction grâce à la sélection de composantes ;
- exclure les composantes qui influencent peu la variable à expliquer ;
- conserver uniquement les composantes qui influencent le plus la variable à expliquer ;
- sélectionner le vrai modèle si celui-ci est présent dans les modèles possibles.

Lorsqu'il y a sélection de modèles, l'inférence est réalisée conditionnellement au modèle sélectionné, et les variables non présentes dans le modèle sont considérées comme non influentes. Le biais dans l'estimation des paramètres les plus influents est généralement faible, mais est possiblement fort pour les variables les moins influentes. Ce phénomène est appelé biais de sélection du modèle. De plus, la sélection de variables conduit à de l'incertitude, et augmente donc la variance des estimateurs.

Nous avons proposé deux estimateurs pour le modèle (1.1). Permettent-ils de réaliser de la sélection de variables automatiquement ?

L'estimateur Ridge conduit à un rétrécissement sans seuillage des coefficients estimés par l'EMCO. Si ceux-ci sont différents de 0, l'estimateur Ridge ne les estimera pas strictement 0.

Pour illustrer ceci, supposons que $\mathbf{X}^T \mathbf{X} = I_d$, alors $\hat{\beta}^{MCO} = \mathbf{X}^T \mathbf{Y}$ et $\hat{\beta}^{Ridge} = \frac{1}{1+\lambda} \mathbf{X}^T \mathbf{Y} = \frac{1}{1+\lambda} \hat{\beta}^{MCO}$. Comme $0 < \frac{1}{1+\lambda} \leq 1$, l'estimateur Ridge ne réalise pas de la sélection automatique des variables. Naturellement, il existe des méthodes permettant de le faire.

1.1.5 Méthodes de sélection de variables fondées sur les tests

Test de nullité d'un coefficient de la régression Si les bruits sont centrés, non corrélés, homoscedastiques et forment un vecteur gaussien, alors pour tout $j \in 1, \dots, d$,

$$\frac{\hat{\beta}_j^{MCO} - \beta_j}{\sqrt{\hat{\sigma}^2 (\mathbf{X}^T \mathbf{X})_{jj}^{-1}}} \sim T(n-d),$$

où $T(n-d)$ est une loi de Student à $n-d$ degrés de liberté et $\hat{\sigma}^2 = \frac{1}{n-d-1} \sum_{i=1}^n (y_i - \hat{y}_i)^2$.

Pour tester la nullité d'un coefficient, c'est-à-dire $(H_0) \beta_j = 0$ versus $(H_1) \beta_j \neq 0$, nous considérons la statistique de test suivante :

$$T(Y) = \frac{\hat{\beta}_j^{MCO}}{\sqrt{\hat{\sigma}^2 (\mathbf{X}^T \mathbf{X})_{jj}^{-1}}},$$

qui suit sous (H_0) la loi $T(n-d)$. Supposons que nous considérons le test à un niveau de confiance α , alors la zone de rejet du test est

$$R_{(H_0)} = \{y \mid |T(y)| \geq t_{n-d}(1 - \alpha/2)\},$$

où $t_{n-d}(1 - \alpha/2)$ est le quantile d'ordre $1 - \alpha/2$ d'une loi de Student à $n-d$ degrés de liberté.

Nous avons exhibé un test de nullité d'une composante. Pour tester simultanément la nullité de toutes les composantes, nous pouvons calculer la significativité de chaque composante et conserver uniquement celles qui sont significatives à un niveau de confiance fixé. Cependant, cette méthode n'est pas optimale, car, lorsque les variables explicatives sont corrélées, la significativité d'une composante dépend de celles des autres composantes.

Méthodes séquentielles de sélection de variables L'une des plus anciennes méthodes de sélection de variables est fondée sur l'application séquentielle de tests d'hypothèse de nullité (voir Efroymson [41] ou Hocking [61]). Cette méthode est particulièrement intuitive. Nous citons l'exemple des méthodes de sélection backward (descendante). Il existe aussi des méthodes de sélection séquentielle forward et stepwise. La méthode backward s'écrit :

Algorithme Sélection backward

1. Calcul du modèle saturé
 2. Test statistique pour connaître les significativités des composantes
 3. S'arrêter si toutes les composantes sont significativement différentes de 0 à un niveau α , sinon supprimer la composante la moins significative
 4. Remplacer le modèle actuel par le modèle privé de la composante supprimée précédemment et recommencer.
-

Supposons que $d = 10$, alors $(10 + 9 + \dots + 1)/(2^{10} - 1) \approx 5\%$ des modèles au maximum sont étudiés. Les méthodes de recherche pas à pas sont fortement dépendantes des données et du point

de départ. En effet, à une itération donnée, nous choisissons le meilleur modèle conditionnellement au modèle sélectionné à l'itération précédente, qui est nécessairement identique ou moins performant que le meilleur modèle à nombre de variables constant. De plus, si deux variables ont une significativité proche, l'une ou l'autre variable peut être exclue du modèle en fonction des données, ce qui influencera la suite de la recherche de modèles. Lorsque les covariables sont corrélées, des variables significatives peuvent apparaître comme non significatives, conduisant à la présence de faux négatifs. La sélection pas à pas avec l'utilisation de tests conduit à un enchaînement de tests, pour lesquels les niveaux de confiance sont incertains. De plus, le critère d'entrée est subjectif : pourquoi choisir un seuil de 0.05 et non de 0.01 ou 0.10 ? Les méthodes de type pas à pas peuvent sembler trop descriptives.

Il existe des critères de sélection de modèles, qui permettent de s'affranchir du niveau de test, subjectif.

1.1.6 Critères de sélection de modèle

Nous utilisons le BIC (Bayesian Information Criterion), l'AIC (Akaike Information Criterion) et le GCV (Generalized Cross Validation), qui sont présentés dans la Table 1.1. Dans cette table, la ligne *Article* correspond à des articles auxquels le lecteur peut se référer pour plus de détails sur la construction des différents critères.

	BIC	AIC	GCV
Article	Schwarz [102]	Akaike [1]	Craven et Wahba [115]
Critère à optimiser	Probabilité a posteriori du modèle sélectionné	Distance de Kullback-Leibler	Distance euclidienne
Ecriture du critère $C(\hat{\beta})$	$\ln(\frac{1}{n}\ \mathbf{Y} - \mathbf{X}\hat{\beta}\ _2^2) + \frac{\ln(n)}{n}\hat{d}f$	$\ln(\frac{1}{n}\ \mathbf{Y} - \mathbf{X}\hat{\beta}\ _2^2) + \frac{2}{n}\hat{d}f$	$\ \mathbf{Y} - \mathbf{X}\hat{\beta}\ _2^2 / \left(n(1 - \frac{\hat{d}f}{n})^2\right)$
Consistence en termes de sélection de modèle	Oui	Non	Non
Vitesse minimax optimale	Non	Oui	Oui

TABLE 1.1 – Critères BIC, AIC et GCV

Le lecteur intéressé par la construction de l'AIC peut notamment se référer à Akaike [1] ou à Anderson et Burnham [23], dont le livre explique l'intérêt de minimiser la pseudo-distance de Kullback-Leibler [76]. Cette pseudo-distance, inspirée par l'entropie de Shannon [103], quantifie l'information perdue lors de l'utilisation de l'estimateur pour approximer la vraie fonction de régression. Le GCV minimise la distance euclidienne. Le BIC, qui est un critère bayésien, sélectionne le modèle le plus probable en fonction des données. Les principes de l'AIC et du GCV sont proches (minimiser un critère quantifiant la différence entre l'estimateur et la vraie fonction). Le BIC a une interprétation différente de l'AIC et du GCV. Ceci explique que ces deux derniers critères aient des propriétés proches et différentes de celles du BIC.

Il existe d'autres critères de sélection de modèle, parmi lesquels le Cp de Mallows¹ (Mallows, [85]), qui minimise l'erreur quadratique moyenne, le critère MDL (Rissanen, [99]), fondé sur la complexité de Kolmogorov [75], et le RIC (Shibata, [104]), qui est une extension de l'AIC.

Ces critères sélectionnent un modèle parmi des modèles candidats, qui peuvent être obtenus avec une recherche exhaustive. Avec 10 variables, il y a $2^{10} - 1 = 1023$ modèles possibles. Lorsque d est élevé, il est donc coûteux de construire et de comparer tous les modèles. Pour une taille de modèle donnée, il est équivalent de minimiser l'AIC, le BIC ou le coefficient de détermination R^2 , où $R^2 = 1 - \frac{\|\mathbf{Y} - \mathbf{X}\hat{\beta}\|_2^2}{\|\mathbf{Y} - \frac{1}{n}\sum \mathbf{Y}_i\|_2^2}$. Pour chaque dimension de modèle, le sous-ensemble de variables avec le

1. $\|\mathbf{Y} - \hat{\beta}\mathbf{X}\|_2^2 / \hat{\sigma}^2 - n + 2d$

R^2 le plus fort peut être identifié, puis comparé avec les autres modèles identifiés de dimension différente. Lorsque d est petit, dans le cas des modèles linéaires gaussiens, nous pouvons utiliser l'algorithme dit de "branch and bound" (voir Furnival et Wilson [48] par exemple) qui permet de déterminer le sous-modèle ayant le R^2 le plus fort pour une dimension de sous-modèle donnée sans nécessiter l'estimation de chaque sous-modèle séparément. Cependant, cet algorithme n'est pas efficace lorsque d est grand ou que les modèles ne sont pas gaussiens. Une stratégie plus structurée est alors nécessaire. Nous pouvons utiliser des méthodes de recherche de type pas à pas, en remplaçant la règle de décision fondée sur les tests par l'AIC, le GCV ou le BIC. Cependant, ces méthodes séquentielles manquent de robustesse.

Nous cherchons des méthodes alternatives, permettant de créer un sous-ensemble de modèles candidats, parmi lesquels nous sélectionnons celui qui minimise un critère de sélection de modèles, type AIC ou GCV (pour les qualités prédictives du modèle sélectionné) ou BIC (pour les qualités sélectives du modèle sélectionné). Les méthodes de régression pénalisée répondent à cet objectif.

1.1.7 Sélection de variables et méthodes de régression pénalisée

1.1.7.1 De l'estimateur Bridge à l'estimateur LASSO

Posons $\lambda \geq 0, \gamma > 0$. L'estimateur Bridge [47] est défini par l'expression suivante :

$$\hat{\boldsymbol{\beta}}^{Bridge} = \arg \min_{\boldsymbol{\beta}} Q^{MCO}(\boldsymbol{\beta}) + \lambda \sum_{i=1}^d |\beta_i|^\gamma.$$

Si $\gamma = 2$, il s'agit de l'estimateur Ridge, qui ne réalise pas de la sélection automatique. Ceci est vérifié dès que $\gamma > 1$.

L'existence de l'estimateur Bridge est assurée par la continuité de la fonction objective associée et par le fait que celle-ci tende vers l'infini lorsque la norme 2 de $\boldsymbol{\beta}$ tend vers l'infini. De plus, lorsque $\gamma = 1$ et que $\underline{\mathbf{X}}$ est injective, la fonction objective est strictement convexe et l'estimateur associé est donc unique. L'estimateur alors défini par (1.7) est l'estimateur LASSO [112].

$$\hat{\boldsymbol{\beta}}^{LASSO} = \arg \min_{\boldsymbol{\beta}} Q^{MCO}(\boldsymbol{\beta}) + \lambda \sum_{i=1}^d |\beta_i|. \quad (1.7)$$

Comme pour l'estimateur Ridge, l'équation (1.7) se ré-écrit sous la forme d'une optimisation sous contrainte :

$$\hat{\boldsymbol{\beta}}^{LASSO} = \arg \min_{\boldsymbol{\beta}} Q^{MCO}(\boldsymbol{\beta}) \text{ s.c. } \sum_{i=1}^d |\beta_i| \leq t. \quad (1.8)$$

Si la matrice $\underline{\mathbf{X}}$ est orthonormale, alors les estimateurs Ridge et LASSO sont équivalents aux transformations de l'estimateur MCO présentées dans la Table 1.2.

Estimateur	Forme
Ridge	$\hat{\boldsymbol{\beta}}^{MCO} / (1 + \lambda)$
LASSO	$sign(\hat{\boldsymbol{\beta}}^{MCO}) \max(0, \hat{\boldsymbol{\beta}}^{MCO} - \lambda)$

TABLE 1.2 – Estimateurs Ridge et LASSO de β_j dans le cas où $\underline{\mathbf{X}}$ est orthonormale

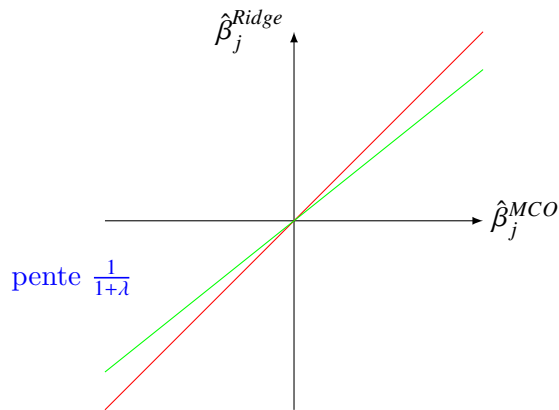


FIGURE 1.2 – Effet de l'estimateur Ridge

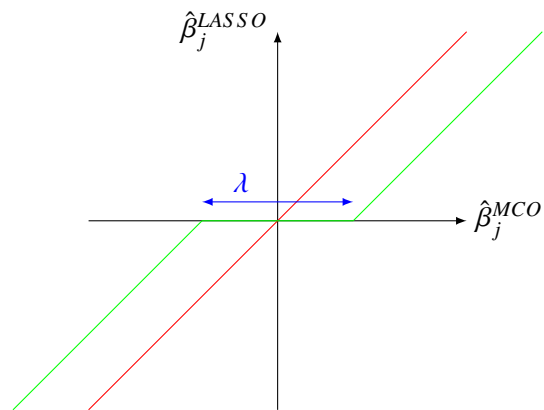


FIGURE 1.3 – Effet de l'estimateur LASSO

sign est l'application qui a un réel positif associe 1 et à un réel négatif -1. Ridge réalise un rétrécissement proportionnel de l'EMCO, alors que le LASSO translate chaque coefficient d'un facteur λ , en tronquant à 0, et réalise donc un rétrécissement et un seuillage.

Considérons par exemple un modèle à deux coefficients, β_1 et β_2 . Les contours de l'erreur quadratique moyenne (EQM) empirique $\|\mathbf{X}\hat{\beta} - \mathbf{X}\beta\|_2^2$ forment des ellipsoïdes centrés sur l'EMCO. La contrainte de la méthode Ridge est constituée du disque délimité par $\beta_1^2 + \beta_2^2 \leq t$ (équation d'un disque) alors que celle du LASSO est délimitée par $|\beta_1| + |\beta_2| \leq t$ (équation d'un carré). L'estimation des coefficients se fait au point où les contours de l'EQM empirique croisent les régions de contrainte. Contrairement au disque, le carré a des angles. Si l'une des ellipsoïdes croise le carré dans un angle, alors l'estimation d'un des paramètres est nulle. La Figure 1.4 illustre la forme des pénalités LASSO et Ridge.

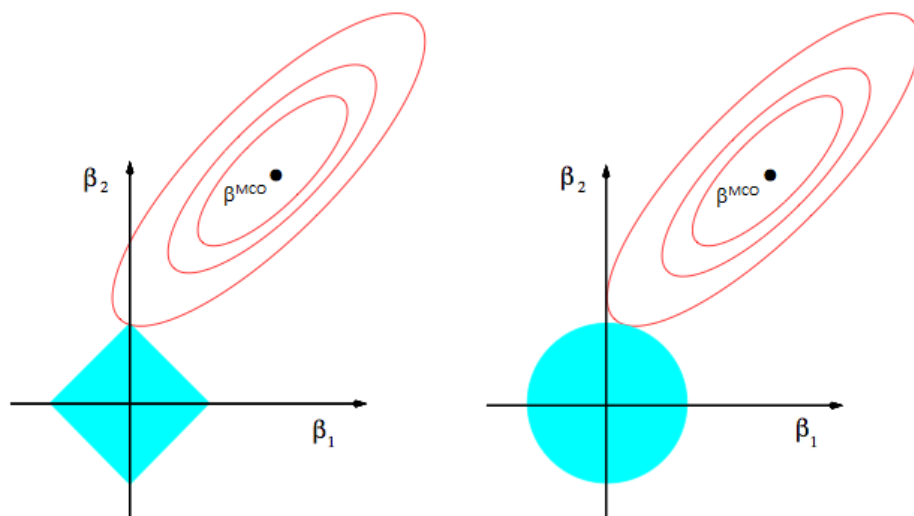


FIGURE 1.4 – Pénalité respectivement LASSO et Ridge lorsqu'il n'y a que deux paramètres à estimer, les domaines bleus correspondant à $p_\lambda(\beta_1, \beta_2) = 1$ et les ellipsoïdes rouges représentant le contour de l'erreur quadratique moyenne (source : Hastie, Tibshirani et Friedman, [58])

1.1.7.2 Nonnegative Garrote

Nous rappelons que nous avons noté $\hat{\beta}^{MCO}$ l'EMCO. $\hat{\beta}^{NNG}(s) = \hat{c}\hat{\beta}^{MCO}$ définit un nouvel estimateur (noté ENNG), avec :

$$\hat{c} \in \arg \min_{c \in \mathbb{R}^d} \|Y - \beta_0 I_n - \sum_{j=1}^d c_j \hat{\beta}_j^{MCO} X_j\|_2^2 \text{ s.c. } c_j \geq 0 \ \forall j \text{ et } \sum_{j=1}^d c_j \leq s.$$

Plus s est faible, plus il y a des c_j nuls et donc de coefficients de $\hat{\beta}^{NNG}(s)$ nuls. Cette procédure, introduite par Breiman [20], est appelée Nonnegative Garrote. La forme lagrangienne du problème d'optimisation est la suivante :

$$\hat{c} \in \arg \min_{c \in \mathbb{R}^d} \|Y - \beta_0 I_n - \sum_{j=1}^d c_j \hat{\beta}_j^{MCO} X_j\|_2^2 + \lambda \sum_{j=1}^d c_j.$$

Dans le cas du plan d'expérience standardisé, l'ENNG est proportionnel à l'EMCO, comme l'estimateur Ridge. Cependant, contrairement à cet estimateur, les poids multipliant l'EMCO associés à l'ENNG peuvent être nuls. L'ENNG permet donc de réaliser de la sélection automatique de variables.

Nous avons présenté deux familles de méthodes de régression pénalisée (Bridge, Nonnegative Garrote). Cependant, la pénalisation peut prendre des formes différentes. Dans le paragraphe suivant, nous donnons les propriétés nécessaires à la pénalisation pour être efficace en termes de sélection de variables et de prédiction.

1.1.7.3 Généralités sur les méthodes de régression pénalisée

Une méthode de Moindres Carrés Pénalisés minimise en fonction de β le critère suivant :

$$Q^{MCO}(\beta) + \sum_{j=1}^d p_\lambda(|\beta_j|),$$

telle que $p_\lambda : \mathbb{R} \mapsto \mathbb{R}^+$ dépend d'un paramètre de lissage λ .

Fan et Antoniadis [7] montrent que p_λ doit posséder les trois propriétés suivantes pour être performante en termes de sélection de variables et d'estimation :

- **Sparsité** : la pénalité doit permettre d'estimer certains coefficients par 0 et donc conduire à des estimateurs creux lorsque le modèle est creux (capacité à réduire le nombre de composantes)
- **Non biaisé** : l'introduction d'une pénalité conduit à un biais sur les estimateurs. Il est souhaité avoir un estimateur de chaque composante j approximativement non biaisé lorsque le coefficient β_j est fort.
- **Continuité** : la pénalité obtenue doit être continue en fonction de l'EMCO.

Contrairement à la pénalité Ridge, la pénalité LASSO remplit l'hypothèse de sparsité. Les estimateurs Ridge et LASSO sont continus en fonction de $\hat{\beta}^{MCO}$. Par contre, lorsque le design est orthonormé, la forme analytique de l'estimateur LASSO correspond à un seuillage doux de l'EMCO, le biais est donc non nul, même pour les coefficients forts.

Pour le modèle linéaire, nous avons à disposition un certain nombre de méthodes performantes permettant de sélectionner les variables et d'estimer les coefficients. Cependant, l'hypothèse de linéarité est très forte.

1.1.8 Limites du modèle linéaire

Sauf pour les phénomènes particulièrement simples, les modèles de régression sont une approximation de la réalité. L'utilisation des modèles linéaires conduit à une modélisation généralement très simplificatrice du phénomène étudié, amenant à un biais de modélisation incompressible.

Considérons la Mean Square Error (MSE) :

$$MSE(\hat{f}) = E\left((\hat{f} - f)^2\right),$$

où f est la vraie fonction de régression et $\hat{f}$ l'estimateur obtenu dans la classe d'estimateurs considérée. Pour la régression linéaire, l'estimateur est de la forme $\hat{f} : x \rightarrow \hat{\beta}^T x$, où $x \in \mathbb{R}^d$ et $\hat{\beta}$ l'estimateur considéré. Alors (voir par exemple Wasserman, [118]), $MSE^{linéaire}(\hat{f}) = V(\hat{f}) + \text{biais}_{linéaire}^2 + O(1/n)$, où $\text{biais}_{linéaire}^2$ est l'erreur minimale qui est commise lorsque la vraie fonction de régression est approchée par une fonction de régression linéaire.

Pour étudier la consommation d'électricité, les modèles linéaires peuvent ne pas être suffisamment souples, car de nombreuses covariables ont des effets non linéaires (voir chapitre 2). Ceci nous conduit à utiliser une autre famille de modèles.

1.2 Généralités sur le modèle additif

1.2.1 Justification du modèle additif

Un modèle de régression peut s'écrire sous la forme suivante :

$$E(Y|X_1 = x_1, \dots, X_d = x_d) = f(x_1, \dots, x_d),$$

ou

$$Y = f(X_1, \dots, X_d) + \varepsilon,$$

où ε bruit *i.i.d.* centré de variance σ^2 et f la fonction de régression inconnue que nous cherchons à estimer. Comme le nombre de données n'est pas infini, nous restreignons la classe de fonctions où f est estimée. Sous la seule hypothèse que $f \in C^{s+1}(\mathbb{R}^d)$, la fonction f peut être estimée avec des méthodes non-paramétriques. Cependant, plus il y a de directions dans lesquelles f doit être estimée, c'est-à-dire plus il y a de covariables, plus la forme de f est difficile à reproduire. D'autre part, plus f est régulière, c'est-à-dire plus s est élevé, plus la fonction de régression est facile à estimer. La forme générale de la MSE d'un estimateur non-paramétrique (voir par exemple Wasserman, [118]) s'écrit sous la forme suivante :

$$MSE^{NP}(\hat{f}) = V(\hat{f}) + \left(E(\hat{f}) - f\right)^2 = V(\hat{f}) + O(n^{-2s/(2s+d)}).$$

L'absence d'hypothèse sur la forme de f conduit à une vitesse de convergence dépendante du nombre de variables. Pour s'affranchir de cette difficulté, nous pouvons poser des hypothèses plus fortes sur la forme de f , ce qui nous ramène au cas des modèles et méthodes paramétriques. Cependant, ces modèles manquent de souplesse pour notre problématique.

Les modèles additifs, introduits par Hastie et Tibshirani [57], réalisent un compromis entre la souplesse des modèles non-paramétriques et la non dépendance de la vitesse de convergence des estimateurs par rapport au nombre de composantes des modèles paramétriques. Le modèle s'écrit :

$$E\left(Y|(X_1, \dots, X_d) = (x_1, \dots, x_d)\right) = \beta_0 + \sum_{i=1}^d f_i(x_i). \quad (1.9)$$

Supposer que f est additive a un coût (introduit un biais) mais cette hypothèse est moins restrictive que supposer la linéarité, donc l'erreur de modélisation est plus faible. Grâce à la propriété d'additivité, il existe des estimateurs du modèle (1.9) pour lesquels la vitesse de convergence des estimateurs est indépendante du nombre de variables :

$$MSE^{additif} = \sigma^2 + \text{biais}_{additif}^2 + O(n^{-(2s/2s+1)}).$$

Indépendamment du nombre de covariables, pour estimer une composante du modèle, nous pouvons considérer toutes les autres connues.

Dans la suite, nous supposons le modèle additif. Son estimation est un problème de dimension infinie et nous approchons ses composantes à l'aide d'une projection dans un sous-espace fini-dimensionnel. Cette approximation influencera l'étude théorique des procédures retenues (voir chapitre 3).

1.2.2 Approximation des composantes dans des bases de B-Splines

1.2.2.1 Définition et propriétés des bases de B-Splines

Les bases de B-Splines ont été introduites par De Boor [33]. Ce sont un assemblage de morceaux de polynômes reliés entre eux. Pour un jeu de noeuds fixé $\xi = (\xi_{-q} = \dots = \xi_{-1} = a = \xi_0, \xi_1, \dots, \xi_K, \xi_{K+1} = b = \xi_{K+2} = \dots = \xi_{K+q+1})$, les B-Splines de degré q sont définies par la formule récursive :

$$\begin{aligned} B_j^0(x) &= \mathbb{1}_{[\xi_j, \xi_{j+1}]}(x), \\ B_j^q(x) &= \frac{x - \xi_j}{\xi_{j+q} - \xi_j} B_j^{q-1}(x) + \frac{\xi_{j+1} - x}{\xi_{j+q+1} - \xi_{j+1}} B_{j+1}^{q-1}(x). \end{aligned}$$

Dans cette définition, $B_j^q(x)$ est la *jème* B-Spline de degré q . Pour construire une base de B-Splines de degré q , il faut au moins $2q + 2$ noeuds. Eilers et Marx [42] affirment les propriétés suivantes pour une B-Spline de degré q :

- une B-Spline comporte $q + 1$ morceaux de polynôme de degré q ,
- pour chaque noeud d'une B-Spline, les dérivées sont continues jusqu'à l'ordre $q - 1$,
- la B-Spline est positive entre ses noeuds inférieur et supérieur, nulle sinon,
- en un point x , il y a $q + 1$ B-Splines non nulles.

Par exemple, l'expression suivante donne la *jème* B-Spline d'une base de B-Splines de degré 2 avec des noeuds uniformément répartis sur un intervalle borné de $\mathbb{R}$:

$$B_j^2(x) = \frac{1}{2}(x - \xi_j)^2 \mathbb{1}_{[\xi_j, \xi_{j+1}]}(x) + \left(\frac{1}{2} + (x - \xi_{j+1}) - (x - \xi_{j+1})^2\right) \mathbb{1}_{[\xi_{j+1}, \xi_{j+2}]}(x) + \frac{1}{2}(1 - (x - \xi_{j+2}))^2 \mathbb{1}_{[\xi_{j+2}, \xi_{j+3}]}(x).$$

Soit m l'ordre des B-Splines, alors $m = K + q + 1$ si une constante est considérée dans la base, $m = K + q$ sinon.

De Boor [33] donne la formule pour calculer la *pième* dérivée pour les fonctions B-Splines :

$$\left(\sum_{j=-q}^K \beta_j B_j^q(x) \right)^{(p)} = \sum_{j=-q+p}^K \beta_j^{(p)} B_j^{q-p}(x),$$

où les coefficients $\beta_j^{(p)}$ sont définis par la formule de récurrence :

$$\begin{aligned}\beta_j^{(1)} &= \frac{q(\beta_j - \beta_{j-1})}{\xi_{j+q} - \xi_j}, \\ \beta_j^{(p)} &= \frac{(q+1-p)(\beta_j^{(p-1)} - \beta_{j-1}^{(p-1)})}{\xi_{j+q+1-p} - \xi_j}.\end{aligned}$$

Les bases de B-Splines ont trois paramètres : le degré des B-Splines, ainsi que le nombre et le placement des noeuds. La bibliographie portant sur le choix de ces deux derniers paramètres est riche. Les noeuds peuvent être placés de manière équidistante ou en fonction d'un certain nombre des quantiles de covariables. Ruppert [100] ou Molinari *et al.* [90] proposent de pénaliser les coefficients des splines. Zhou et Shen [123] ou Gervini [52] utilisent des méthodes de sélection pas à pas de noeuds.

Dans le paragraphe suivant, nous présentons une méthode pour placer les noeuds dans le cas univarié, puis introduisons un algorithme de type backfitting pour généraliser la méthode au cas du modèle additif. Cette méthode illustre le caractère local des B-Splines, justifiant ainsi leur utilisation.

1.2.2.2 Un exemple de placement des noeuds : méthode de Zhou et Shen [123]

Zhou et Shen [123] proposent une méthode adaptative de placement des noeuds dans le cas de modèle univarié. Le modèle considéré est $y_i = f(x_i) + \varepsilon_i$, sous les conditions usuelles et avec $f \in C^{s+1}(\mathbb{R})$.

La méthode consiste en une sélection stepwise des noeuds. Le risque de Stein (1.10) est utilisé pour sélectionner les modèles.

$$R(\hat{f}) = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{f}(x_i; \xi))^2 + C \frac{(k+q+1)\sigma^2}{n}, \quad (1.10)$$

où k le nombre de noeuds internes et σ estimée par l'estimateur robuste $\hat{\sigma} = \text{mediane}\left(\frac{|y_{2i} - y_{2i-1}|}{0.6745\sqrt{2}}, i = 1, \dots, n/2\right)$. Les auteurs suggèrent de considérer $C = 2$. Notons $\mathbf{B}_\xi(-)$ la base de B-Splines de noeuds ξ , $\mathbf{B}_\xi = (\mathbf{B}_\xi(x_1)^T, \dots, \mathbf{B}_\xi(x_n)^T)^T$, $\mathbf{y} = (y_1, \dots, y_n)$ et $\hat{f}(x; \xi) = \mathbf{B}_\xi(\mathbf{x})(\mathbf{B}_\xi^T \mathbf{B}_\xi)^{-1} \mathbf{B}_\xi^T \mathbf{y}$.

Le calcul du risque de Stein est coûteux informatiquement, contrairement à celui de la différence entre deux risques de Stein. En effet, grâce aux propriétés d'orthogonalité des moindres carrés, $R(\hat{f}(\cdot, \xi)) - R(\hat{f}(\cdot, \xi^*)) = \sum_{i=1}^n (\hat{f}(x_i; \xi^*) - \hat{f}(x_i; \xi))^2 - 2\sigma^2/n$, avec $\xi^* = \xi \cup \xi^*$. Le premier théorème de Zhou et Shen [123] montre que l'impact de l'ajout d'un noeud est faible lorsqu'il y a beaucoup de noeuds entre x et le noeud ajouté, restreignant ainsi le nombre d'observations à utiliser lors de l'évaluation de la différence de deux risques de Stein. Le second théorème implique que les estimateurs peuvent être étudiés localement, diminuant ainsi le nombre de noeuds autour du noeud à ajouter. Ces deux théorèmes réduisent les temps de calculs grâce à l'approximation suivante :

$$R(\hat{f}(\cdot, \xi)) - R(\hat{f}(\cdot, \xi^*)) \approx \frac{1}{n} \sum_{x_i \in [\xi_{i^*} - M, \xi_{i^*} + M + 1]} (\hat{f}(x_i; \xi_{i^*}) - \hat{f}(x_i; \xi^*))^2 - \frac{2\sigma^2}{n}, \quad (1.11)$$

où M telle que $M \ll k$, $\xi_{i^*} = (\xi_{i^*-M}, \dots, \xi_{i^*+M+1})$ et $\xi_{i^*} = (\xi_{i^*-M}, \dots, \xi_{i^*}, \xi^*, \xi_{i^*+1}, \dots, \xi_{i^*+M+1})$.

L'approximation faite sur la différence entre deux risques de Stein illustre un avantage important des bases de B-Splines qui est de travailler localement. L'algorithme de [123] comporte deux phases principales qui se répètent jusqu'à convergence : une phase d'ajout des noeuds et une phase de suppression et de relocalisation des noeuds, et utilise la différence de deux risques de Stein.

Algorithme SARS de Zhou et Shen [123]

1. **Initialisation** : Connaissances métiers, noeuds équi-distant,...
2. **Recherche des noeuds candidats** : Soit $\xi = (a = \xi_0 < \xi_1, \dots, < \xi_k < \xi_{k+1} = b)$ la séquence de noeuds courante. S'il s'agit des noeuds initiaux, pour tout $i \in \llbracket 0, k \rrbracket$, on cherche un noeud candidat dans $[\xi_i, \xi_{i+1}]$. Sinon, chercher un noeud candidat dans $[\xi_i, \xi_{i+1}]$ si au moins un noeud parmi $\xi_{i-2}, \dots, \xi_{i+3}$ a été ajouté lors de l'étape d'addition de noeuds précédente. Notons $\phi \subset \xi$ la liste des noeuds ajoutés précédemment ($\phi = \xi$ la première fois). L'étape de recherche des noeuds s'écrit :
 Noter $\Xi = \xi$. Pour i allant de 0 à k , si il y a $\xi_j \in \phi$ tel que $|i - j| \leq 2$, alors résoudre avec l'approximation (1.11) le problème d'optimisation :

$$\max_{\xi^* \in]\xi_i, \xi_{i+1}[} (R(\hat{f}(\cdot, \xi)) - R(\hat{f}(\cdot, \xi^*))),$$

puis ajouter le noeud potentiel à Ξ et à ϕ si la différence est positive.

$\xi = \Xi$.

3. **Suppression et relocalisation des noeuds** : Suite à l'étape précédente, on a k_a noeuds internes. Notons $\Xi = \xi$. Pour i allant de 0 à k_a faire, s'il y a $\xi_j \in \phi$ tel que $|i - j| = 1$, alors supprimer ξ_i et résoudre le problème d'optimisation grâce à (1.11) :

$$\max_{\xi^* \in]\xi_{i-1}, \xi_{i+1}[} (R(\hat{f}(\cdot, \xi_{\text{privé de } \xi_i})) - R(\hat{f}(\cdot, \xi^*))),$$

puis ajouter le noeud potentiel à Ξ et à ϕ si la différence est positive.

4. **Fin de l'algorithme** : La première fois qu'aucun noeud n'est ajouté à l'étape 2, effectuer l'étape 3 en considérant les noeuds comme initiaux et ϕ les noeuds ajoutés. La seconde fois s'arrêter.
-

Ensuite, cet algorithme a été intégré dans un algorithme de type backfitting pour être généralisé au modèle additif. Par souci d'identifiabilité, nous supposons que $\sum_{i=1}^n f_j(X_{j,i}) = 0$ pour tout $j \in \llbracket 1, d \rrbracket$. Hastie et Tibshirani [57] ont introduit l'algorithme de backfitting pour estimer les composantes du modèle additif. Dans le backfitting, nous construisons S à l'aide de l'algorithme SARS [57]. Le backfitting se ramène à chaque étape à un problème univarié.

L'application pratique de cette méthode a été insatisfaisante. Nous avons testé cette procédure sur des postes sources (voir l'introduction ou le chapitre 2), qui sont des données de consommation électrique locale relativement désagrégées. Nous avons utilisé la méthode en utilisant un modèle expliquant la consommation par la température et par le moment dans l'année. En plus des défauts propres au backfitting (lenteur de la convergence, problème lorsque les variables sont corrélées), les données sont sur-appprises : un nombre excessif de noeuds est utilisé. Pour corriger ce défaut, nous avons envisagé deux solutions, qui consistent soit à optimiser le risque de Stein, ce qui est coûteux informatiquement, soit à ajouter un terme de lissage à la dernière étape du backfitting. La seconde solution est moins coûteuse parce qu'il n'y a pas de retour sur la base de B-Splines, les effets étant juste lissés. Nous n'avons pas testé la méthode, car nous nous attendions à un faible gain en performance de prédiction.

Algorithme Backfitting de Hastie et Tibshirani [57]

1. **Initialisation** : $\beta_0 = 1/n \sum_{i=1}^n Y_i$ et $f_j^{(1)} \equiv 0$ pour tout $i, j, m=1$,
 2. **Boucle while** : Jusqu'à ce que $RSS_m = \sum_{i=1}^n (Y_i - \beta_0 - \sum_{j=1}^p f_j^{(m)}(X_{j,i}))^2$ arrête de décroître ou vérifie un critère d'arrêt :
 - (a) $m = m + 1$
 - (b) Pour $j \in \llbracket 1, d \rrbracket$
 - $\mathbf{R}_j^{(m)} = (R_{j,i}^{(m)})$ avec $R_{j,i}^{(m)} = Y_i - \beta_0 - \sum_{k \in \llbracket 1, d \rrbracket \setminus j} f_k^{(m)}(X_{ki})$,
 - $f_j^{(m)} = S[\mathbf{R}_j^{(m)}]$, (S matrice de lissage-projection construite à l'aide de SARS)
 - $f_j^{(m)} = f_j^{(m)} - 1/n \sum_{i=1}^n f_j^{(m)}(X_{j,i})$ (pour l'identifiabilité, effets centrés).
-

Notre principal objectif consiste à sélectionner des variables et donc à supprimer des composantes globales. Chercher à résoudre simultanément le problème du nombre de composantes ainsi que du nombre et de la position des noeuds pour chaque composante est un problème trop complexe. Finalement, nous avons choisi de placer les noeuds des bases de B-Splines en fonction d'un certain nombre de quantiles des covariables, pour nous assurer d'avoir des observations entre chaque noeud, ou le long d'une grille uniforme, pour faciliter l'obtention de résultats théoriques. Ensuite, nous utilisons des méthodes de régularisation, type P-Splines, pour estimer des effets lisses.

Dans la sous-section suivante, nous écrivons le modèle (1.9) approximé dans des bases de B-Splines.

1.2.3 Ecriture du modèle additif avec des composantes projetées dans des bases de B-Splines

Les variables explicatives sont projetées dans des bases de B-Splines pour approximer les composantes. Nous pouvons écrire un modèle (1.12) approximant le modèle (1.9).

$$E(Y|X_1, \dots, X_d = (x_1, \dots, x_d)) = \beta_0 + \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} \mathbf{B}_{j,k}^{q_j}(x_j), \quad (1.12)$$

avec $\mathbf{B}_j(-) = (\mathbf{B}_{j,k}^{q_j}(-) | k = 1, \dots, K_j + q_j = m_j)$ la base de B-Splines de degré q_j et K_j noeuds, où est projetée la variable X_j . Soit $\mathbf{B}_j = (\mathbf{B}_j(x_{1j})^T, \dots, \mathbf{B}_j(x_{nj})^T)^T \in \mathbb{R}^{n \times m_j}$. Les paramètres à estimer du modèle sont alors β_0 et $\boldsymbol{\beta} = (\beta_{1,1}, \dots, \beta_{d,m_d})^T$. Les variables explicatives sont les $(X_j)_{j \in \llbracket 1, d \rrbracket}$. Elles peuvent être influentes, ou non, sur la réponse Y . $x_{i,j}$ constitue la *ième* observation de la variable X_j et les paramètres à estimer du modèle sont β_0 et $\boldsymbol{\beta} = (\beta_{1,1}, \dots, \beta_{d,m_d})^T$. Les $(m_i)_{i \in \llbracket 1, d \rrbracket}$ sont connus. Sauf contradiction, $m_i = m = m_n$ pour tout i .

Le modèle (1.12) est une approximation du modèle (1.9). Le caractère local des B-Splines permet d'avoir une bonne approximation, qui nous ramène au cas d'un modèle linéaire. Même si le nombre de variables est inférieur au nombre d'observations, cette approximation peut nous conduire au cas de la grande dimension puisque le nombre de paramètres à estimer est $dm_n + 1$.

1.2.4 Calcul du degré de liberté

Puisque nous nous sommes ramenés au cas d'un modèle linéaire, nous avons un estimateur de la forme $\mathbf{S}_\lambda \mathbf{Y} = \mathbf{B}^T \hat{\boldsymbol{\beta}}$, où $\mathbf{B} = (\mathbf{B}_1, \dots, \mathbf{B}_d)$ et $\mathbf{S}_\lambda = \mathbf{B}(\mathbf{B}^T \mathbf{B})^{-1} \mathbf{B}^T$. Le nombre de degrés de liberté est $n - \text{trace}(\mathbf{S}_\lambda)$, ce qui revient à compter le nombre de composantes non nulles.

Souvent, certaines composantes du modèle (1.9) sont nulles. Par souci d'écriture, seules les D premières composantes sont supposées non nulles. La section suivante établit des procédures de sélection et d'estimation pour le modèle (1.12) en combinant des méthodes de régression pénalisée qui peuvent à la fois sélectionner les variables et estimer les effets. Il y a une notion de groupe à considérer. En effet, il est associé à chaque variable un groupe de coefficients, parce que chaque covariable est projetée dans une base de B-Splines.

1.3 Procédures de sélection et d'estimation

1.3.1 Méthode de régression pénalisée

1.3.1.1 Group LASSO

Pour les modèles linéaires, le LASSO permet de sélectionner les variables. Nous utilisons une version groupée du LASSO pour sélectionner les composantes du modèle additif approximées dans des bases de B-Splines.

Version non standardisée Le Group LASSO a été introduit par Yuan et Lin [122]. Notons $\beta_j = (\beta_{j1}, \dots, \beta_{jm_j})$ et m_j l'ordre de la base de B-Splines où est projetée la *jème* variable. Appliquée au modèle (1.12), la fonction objective du Group LASSO est donnée par l'équation (1.13).

$$Q^{\text{GroupLASSO}}(\beta) = Q^{\text{MCO}}(\beta) + \lambda \sum_{j=1}^d \sqrt{m_j} \|\beta_j\|_2. \quad (1.13)$$

Il y a d groupes pénalisés séparément, constitués chacun d'un groupe de coefficients associés à la projection d'une variable dans une base de B-Splines. La constante n'est pas pénalisée.

L'estimateur Group LASSO minimise la fonction objective (1.13) :

$$\hat{\beta}^{\text{GroupLASSO}} = \arg \min Q^{\text{GroupLASSO}}(\beta).$$

Comme le LASSO, le Group LASSO permet de sélectionner les composantes. Il réalise de la sélection inter-groupe, mais pas intra-groupe.

Plusieurs algorithmes minimisant l'équation (1.13) utilisent un plan d'expérience dans lequel les groupes sont orthonormalisés. Si ce n'est pas le cas à l'origine, il existe des méthodes permettant de les orthonormaliser. Cependant, si nous ré-exprimons le critère (1.13), nous changeons le problème. En effet, si $\mathbf{B}_j$ est de plein rang, nous pouvons l'écrire $\mathbf{B}_j = \Psi_j \Phi_j$, avec Ψ_j une matrice avec des colonnes orthonormales de dimension $n \times m_j$ et Φ_j une matrice inversible de dimension $m_j \times m_j$. Alors, l'expression (1.13) devient

$$Q^{\text{GroupLASSO}}(\beta) = \|Y - \sum_{j=1}^d \Psi_j \Phi_j \beta_j\|_2^2 + \lambda \sum_{j=1}^d \sqrt{m_j} \|\beta_j\|_2.$$

Posons $\theta_j = \Phi_j \beta_j$. Alors le critère s'écrit :

$$Q^{\text{GroupLASSO}}(\theta) = \|Y - \sum_{j=1}^d \Psi_j \theta_j\|_2^2 + \lambda \sum_{j=1}^d \sqrt{m_j} \|\Phi_j^{-1} \theta_j\|_2.$$

Le critère ainsi ré-écrit est différent de $\|Y - \sum_{j=1}^d \Psi_j \theta_j\|_2^2 + \lambda \sum_{j=1}^d \sqrt{m_j} \|\theta_j\|_2$, qui est le critère minimisé par les algorithmes travaillant avec un plan d'expérience dont les groupes sont orthonormalisés.

Version standardisée L'expression (1.13) ne tient pas compte des échelles. En effet, les pénalisations ne varient d'un groupe à l'autre qu'en fonction de la dimension du groupe. En standardisant chaque groupe, nous homogénéisons les variabilités des variables. Puisque toutes les variables ont alors la même échelle, la forme de la pénalisation devient légitime. De plus, en supprimant les effets d'échelle, nous rendons possible la comparaison de l'influence d'une variable par rapport à une autre. La standardisation permet de gagner en interprétabilité.

La version standardisée du Group LASSO consiste à considérer le problème (1.14) :

$$Q^{GroupLASSO2}(\boldsymbol{\beta}) = Q^{MCO}(\boldsymbol{\beta}) + \lambda \sum_{j=1}^d \sqrt{m_j} \|\mathbf{B}_j \boldsymbol{\beta}_j\|_2. \quad (1.14)$$

Si $\mathbf{B}_j$ est de plein rang, nous pouvons l'écrire $\mathbf{B}_j = \Psi_j \Phi_j$, avec Ψ_j une matrice avec des colonnes orthonormales de dimension $n \times m_j$ et Φ_j une matrice inversible de dimension $m_j \times m_j$. De plus, posons $\boldsymbol{\theta}_j = \Phi_j \boldsymbol{\beta}_j$. Alors, l'expression (1.14) devient

$$Q^{GroupLASSO2}(\boldsymbol{\beta}) = \|Y - \sum_{j=1}^d \Psi_j \boldsymbol{\theta}_j\|_2^2 + \lambda \sum_{j=1}^d \sqrt{m_j} \|\boldsymbol{\theta}_j\|_2.$$

Les colonnes de Ψ_j sont orthonormales. L'expression (1.14) se ré-écrit par l'expression (1.15) :

$$Q^{GroupLASSO2}(\boldsymbol{\theta}) = \|Y - \sum_{j=1}^d \Psi_j \boldsymbol{\theta}_j\|_2^2 + \lambda \sum_{j=1}^d \sqrt{m_j} \|\boldsymbol{\theta}_j\|_2. \quad (1.15)$$

La solution du Group LASSO ne peut pas toujours être calculée analytiquement et des méthodes numériques sont nécessaires pour l'évaluer.

Algorithme utilisé L'algorithme *Group Coordinate Descent*, proposé notamment par Huang *et al.* [66], a été utilisé. Cet algorithme découle d'une approximation du premier ordre de Taylor de la fonction objective. Il faut au préalable orthogonaliser la matrice de design. Pour cela, nous pouvons utiliser une décomposition de Cholesky pour les d matrices de Gram des groupes : $\frac{1}{n} \mathbf{B}_j^T \mathbf{B}_j = \mathbf{U}_j^T \mathbf{U}_j$, avec $\mathbf{U}_j$ une matrice triangulaire supérieure.

Soit $\tilde{\mathbf{B}}_j = \mathbf{B}_j \mathbf{U}_j^{-1}$. Alors $\frac{1}{n} \tilde{\mathbf{B}}_j^T \tilde{\mathbf{B}}_j = \frac{1}{n} \mathbf{U}_j^{-T} \mathbf{B}_j^T \mathbf{B}_j \mathbf{U}_j^{-1} = \frac{1}{n} \mathbf{U}_j^{-T} \mathbf{U}_j^T \mathbf{U}_j \mathbf{U}_j^{-1} = \mathbf{I}$ et $\tilde{\mathbf{B}}_j \tilde{\boldsymbol{\beta}}_j = \tilde{\mathbf{B}}_j \mathbf{U}_j \mathbf{U}_j^{-1} \tilde{\boldsymbol{\beta}}_j = \mathbf{B}_j (\mathbf{U}_j^{-1} \tilde{\boldsymbol{\beta}}_j)$, et finalement $\boldsymbol{\beta}_j = \mathbf{U}_j^{-1} \tilde{\boldsymbol{\beta}}_j$. Il est supposé pour l'écriture de l'algorithme suivant que les $\mathbf{B}_j$ sont orthogonales.

Algorithme Group Coordinate Descent de Brehenny *et al.* [19]

$\hat{\boldsymbol{\beta}}^{(0)} = (\boldsymbol{\beta}_1^{(0)T}, \dots, \boldsymbol{\beta}_d^{(0)T})^T$ valeurs initiales.

1. **Initialisation** : $i = 0$ et $\mathbf{r} = \mathbf{Y} - \sum_{j=1}^d \mathbf{B}_j \boldsymbol{\beta}_j^{(0)T}$
 2. **Boucle Pour** : Pour $j \in \llbracket 1, d \rrbracket$
 - (a) Calculer $\mathbf{z}_j = \frac{1}{n} \mathbf{B}_j^T \mathbf{r} + \hat{\boldsymbol{\beta}}_j^{(i)}$
 - (b) Calculer $\hat{\boldsymbol{\beta}}_j^{(i+1)} = \left(1 - \frac{\lambda \sqrt{m_j}}{\|\mathbf{z}_j\|_2}\right)_+ \mathbf{z}_j$
 - (c) Calculer $\mathbf{r} = \mathbf{r} - \mathbf{B}_j (\hat{\boldsymbol{\beta}}_j^{(i+1)} - \hat{\boldsymbol{\beta}}_j^{(i)})$
 3. $i = i + 1$
 4. Répéter 2-3 jusqu'à convergence
-

Pourquoi utiliser la forme de l'étape 2.(a) de l'algorithme ?

$$\begin{aligned}
\frac{1}{n} \mathbf{B}_j^T \mathbf{r} + \hat{\boldsymbol{\beta}}^{(s)} &= \frac{1}{n} \mathbf{B}_j^T (\mathbf{Y} - \sum_{i=1}^d \mathbf{B}_i \boldsymbol{\beta}_i^{(s)T}) + \hat{\boldsymbol{\beta}}^{(s)} \\
&= \frac{1}{n} \mathbf{B}_j^T (\mathbf{Y} - \sum_{i=1, j \neq i}^d \mathbf{B}_i \boldsymbol{\beta}_i^{(s)T}) - \hat{\boldsymbol{\beta}}^{(s)} + \hat{\boldsymbol{\beta}}^{(s)} \\
&= \frac{1}{n} \mathbf{B}_j^T (\mathbf{Y} - \sum_{i=1, j \neq i}^d \mathbf{B}_i \boldsymbol{\beta}_i^{(s)T}).
\end{aligned}$$

La forme de l'étape 2.(a) est choisie plutôt que le dernier membre de droite de l'expression précédente, parce qu'elle ne nécessite que $2n$ calculs au lieu des nd nécessaires à la formulation directe des résidus. L'étape 2.(b) correspond à un seuillage faible. L'algorithme a deux avantages : chaque calcul est rapide et l'algorithme est stable : soit la fonction à optimiser décroît, soit elle ne change pas. Il n'est valable que pour un paramètre de lissage λ , dont le choix est compliqué. Nous pouvons répéter cet algorithme sur une grille de λ . Au début, $\lambda = \lambda_{\max} = \max_j \left(\left\| \frac{1}{n} \frac{\mathbf{B}_j \mathbf{Y}}{\sqrt{m_j}} \right\|_2 \right)$ seuille tous les coefficients à 0. Ensuite, on procède le long de la grille en considérant que les coefficients du point précédant de la grille sont les coefficients initiaux de l'algorithme suivant.

Convergence de l'algorithme Tseng [113] donne les résultats permettant de conclure à la convergence de l'algorithme du *Group Coordinate Descent*.

L'objectif de Tseng est de minimiser avec cet algorithme une fonction du type :

$$f(x_1, \dots, x_d) = f_0(x_1, \dots, x_d) + \sum_{k=1}^d f_k(x_k),$$

avec $x_i \in \mathbb{R}^{m_i}$. Nous définissons les termes suivants :

- $\text{dom}(f) = \{x \in \mathbb{R}^{\sum m_i} | f(x) < \infty\}$
- f quasi-convexe si $f(\lambda x + (1 - \lambda)y) \leq \max(f(x), f(y))$
- Nous dirons que f est hemivariable si f n'est constante sur aucun vecteur appartenant à $\text{dom}(f)$. Un exemple de fonction non hemivariable est une fonction f telle qu'il existe $(x_1, y_1) \in A_1 \times A_2$ et $(x_2, y_2) \in A_1 \times A_2$ (avec A_1 et A_2 ensembles quelconques) tels que pour tout $t \in [0, 1]$, $f(tx_1 + (1 - t)x_2, ty_1 + (1 - t)y_2) = \text{constante}$.

Soient les conditions suivantes :

- A1 f_0 continue sur $\text{dom}(f)$
- A2 Pour tout $k \in \llbracket 1, d \rrbracket$ et $(x_j)_{j \neq k}$, $x_k \mapsto f(x_1, \dots, x_k, \dots, x_d)$ est quasi-convexe et hemivariable
- A3 $f_0, \dots, f_d$ continues à gauche
- A4 $\text{dom}(f_0)$ est un ouvert et f_0 tend vers l'infini en chaque point à la frontière de $\text{dom}(f_0)$
- A5 $\text{dom}(f_0) = A_1 \times \dots \times A_d$, pour $A_k \subseteq \mathbb{R}^{m_k}$, avec $k \in \llbracket 1, d \rrbracket$

Sous les conditions A1, A2, A3 et A4 ou A5, Tseng donne des résultats permettant d'affirmer la convergence de l'algorithme. Pour le Group LASSO, $f_0(\boldsymbol{\beta}) = \|\mathbf{Y} - \mathbf{B}\boldsymbol{\beta}\|_2^2$ et pour tout k , $f_k(\boldsymbol{\beta}_k) = \lambda \sqrt{m_k} \|\boldsymbol{\beta}_k\|$, les hypothèses A1, A2, A3, A5 sont vérifiées, donc l'algorithme converge dans le cas du Group LASSO. De plus, comme la fonction objective du Group LASSO est convexe, le minimum est un minimum global (pas nécessairement unique).

Il existe d'autres algorithmes permettant de minimiser la fonction objective du Group LASSO. Par exemple, Breheny et Huang [19] proposent l'algorithme dit de *Local Coordinate Descent*, qui ne demande pas à ce que les groupes de la matrice soient orthonormalisés et met à jour coefficients par coefficients. Meier, Van de Geer et Bühlmann [89] proposent un autre algorithme.

Nous avons choisi l'algorithme de *Group Coordinate Descent* par rapport aux deux autres pour des raisons différentes. Breheny et Huang [19] mettent à jour coefficient par coefficient,

contrairement aux deux autres algorithmes qui mettent à jour groupe de coefficients par groupe de coefficients. L'algorithme de Meier *et al.* [89], codé dans le package R **grplasso**, a un programme d'optimisation à l'intérieur d'une boucle *for*, contrairement à l'algorithme proposé par Huang *et al.* [66], ce qui explique notre choix d'algorithme. Il est codé dans le package R **grpreg**.

Bach [13] établit des conditions sous lesquelles le Group LASSO est consistant en sélection. Le Group LASSO permet de sélectionner des composantes. Nous avons à disposition des algorithmes qui convergent et qui permettent de résoudre le problème de Group LASSO. Cependant, comme le LASSO, l'estimateur Group LASSO introduit un biais dans l'estimation, ce qui nous amène à étudier d'autres méthodes de régression pénalisée.

1.3.1.2 Penalized Spline (P-Splines)

P-Splines d'Eilers et Marx Nous utilisons des P-Splines, introduites par Eilers et Marx [42]. La fonction objective à optimiser est donnée dans l'équation (1.16) :

$$Q^{psplines}(\beta) = Q^{MCO}(\beta) + \sum_{j=1}^d \lambda_j \|D_{2,j}\beta_j\|_2^2, \quad (1.16)$$

où $D_{2,j}$ est la représentation matricielle de la différence d'ordre 2. Si $j > 2$, nous avons $D_{2,j}\beta_j = \beta_j - 2\beta_{j-1} + \beta_{j-2}$.

Lorsque les noeuds sont équi-répartis, cela revient à minimiser l'expression suivante :

$$Q^{psplines}(\beta) = Q^{MCO}(\beta) + \lambda \int_0^1 \left\{ \left(\sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(x_j) \right)^{(2)} \right\}^2 dx,$$

Les dérivées secondes des composantes sont pénalisées. Les effets estimés sont donc lissés.

Wood [119] remarque que :

$$Q^{pspline}(\beta) = \left\| \begin{pmatrix} \mathbf{Y} \\ \mathbf{0} \end{pmatrix} - \begin{pmatrix} \mathbf{B}_0 \\ \Omega_\lambda \end{pmatrix} \Gamma^T \right\|_2^2,$$

$$\text{où } \mathbf{B}_0 = (\mathbb{1}_n, \mathbf{B}), \Gamma = (\beta_0, \beta) \text{ et } \Omega_\lambda = \begin{pmatrix} 0 & \dots & 0 \\ 0 & \lambda_1 D_{k,1} & 0 & \dots & 0 \\ 0 & 0 & \ddots & \dots & 0 \\ 0 & \dots & \dots & \lambda_d D_{k,d} \end{pmatrix},$$

alors

$$\hat{\beta}^{psplines} = (\mathbf{B}_0^T \mathbf{B}_0 + \Omega_\lambda^T \Omega_\lambda)^{-1} \mathbf{B}_0^T \mathbf{Y}.$$

Modification de la pénalité P-Splines Les P-Splines lissent les estimateurs, diminuant ainsi la variance, mais ne permettent pas de réaliser de la sélection. Cette pénalité ne tient pas compte du lissage des fonctions qui sont dans le noyau de la pénalité. Marra et Wood [86] proposent deux modifications de la pénalité associée aux P-Splines. Pour la composante j , la pénalité s'écrit $S_j = D_{2,j} D_{2,j}^T$, avec $D_{2,j}$ de dimension $m_j \times (m_j - 2)$, avec $2 < m_j$. La pénalité S_j peut être décomposée telle que $S_j = U_j \Lambda_j U_j^T$, où U_j et Λ_j matrices respectivement de vecteurs propres et diagonale avec pour valeurs les valeurs propres. Par construction de la matrice de pénalité, il existe des valeurs propres nulles. Marra et Wood [86] proposent alors de modifier la pénalité pour ne pas avoir de valeur propre nulle.

La modification la plus simple consiste à remplacer les valeurs nulles de la diagonale de Λ_j par une petite proportion de la plus petite des valeurs propres strictement positives. Si une telle matrice

est notée $\tilde{\Lambda}_j$, alors la nouvelle pénalité considérée est $\tilde{S}_j = U_j \tilde{\Lambda}_j U_j^T$. La proportion est choisie arbitrairement.

L'autre transformation de la pénalité consiste à ajouter une seconde pénalité $S_j^{(2)} = U_j^* U_j^{*T}$, où U_j^* est la matrice des vecteurs propres associés aux valeurs nulles. La composante j est alors pénalisée par $\lambda_j S_j + \lambda_j^* S_j^{(2)}$. La pénalité ainsi construite est de plein rang.

Propriétés asymptotiques dans le cas univarié Pour ce paragraphe, $d = 1$. Le problème considéré est l'estimation de la fonction f dans le modèle :

$$Y_i = f(X_i) + \varepsilon_i, \quad i = 1, \dots, n,$$

où n le nombre d'observations, $(Y_i)_{i \in \llbracket 1, n \rrbracket}$ la variable à expliquer, $(X_i)_{i \in \llbracket 1, n \rrbracket}$ la variable explicative, $E(\varepsilon_i) = 0$ et $E(\varepsilon_i^2) = \sigma^2$. Il est supposé de plus que pour tout i , si x_i réalisation de X_i , $x_i \in [0, 1]$, $f \in C^q[0, 1]$.

Nous notons $\hat{f}^{PSP}$ l'estimateur P-Splines de f univariée. Claeskens *et al.* [31] ont étudié les propriétés asymptotiques de cet estimateur.

Théorème 1.1 Soit $\tilde{c}_1$ la constante introduite dans le Lemme 3 de Claeskens *et al.* [31]. Soit $K_p = (m - p)(\lambda \tilde{c}_1)^{1/(2p)} n^{-1/(2p)}$. Sous les conditions (A1), (A2) et (A3) de l'article,

1. Si $K_p < 1$, $f \in C^{q+1}[0, 1]$, $K \sim C_1(n^{1/(2q+3)})$ avec C_1 constante et $\lambda = O(n^\gamma)$ avec $\gamma \leq (q + 2 - p)/(2q + 3)$, alors

$$AMSE \triangleq \frac{1}{n} E \left((\hat{f}^{PSP} - f)^T (\hat{f}^{PSP} - f) \right) = O(n^{-(2q+2)/(2q+3)}).$$

2. Si $K_p > 1$, $f \in W^q[0, 1]$, $K \sim C_2(n^\nu)$ avec C_2 constante et $\nu \geq 1/(2q + 1)$ et $\lambda = O(n^{1/(2q+1)})$ tel que $\lambda n^{2q-1} \mapsto \infty$, alors

$$AMSE = O(n^{-(2q)/(2q+1)}).$$

Propriétés asymptotiques dans le cas multivarié Nous nous replaçons dans le cadre du modèle additif (1.9) tel que les composantes soient centrées et de degré de lissage s . Antoniadis, Gijbels et Verhasselt [9] ont étudié l'estimateur P-Splines combiné au backfitting. Sous un certain nombre de conditions, posées notamment sur le nombre de noeuds, sur la distribution des covariables, sur l'existence des dérivées premières des composantes carrés intégrables et sur le caractère lipschitzien de chaque composante, Antoniadis *et al.* [9] établissent que :

Théorème 1.2

$$1/n \|\hat{f}_j^{PSPad} - f_j\|_2^2 = O_P(n^{-(2s+2)/(2s+3)}),$$

en notant $\hat{f}^{PSPad}$ l'estimateur obtenu à l'aide de P-Splines dans le cas additif.

Grâce à la propriété d'additivité, la vitesse de convergence optimale pour un modèle non paramétrique est obtenue pour l'estimateur P-Splines combiné au backfitting dans le cas de modèle additif.

Antoniadis *et al.* [9] montrent que lorsque les P-Splines sont utilisées comme estimateur initial, le Nonnegative Garrote est consistant en sélection.

Sous certaines conditions, le Group LASSO est consistant en sélection (voir Bach [13]). Cependant, comme le LASSO, cette méthode introduit un biais, que nous cherchons à corriger. Sous d'autres conditions, les P-Splines convergent en moyenne quadratique. Nous souhaitons combiner les deux méthodes de régression pénalisée pour obtenir des propriétés simultanées de sélection et d'estimation. Belloni *et al.* [16] ont fait un travail similaire dans le cas des modèles linéaires, en combinant le LASSO avec le MCO.

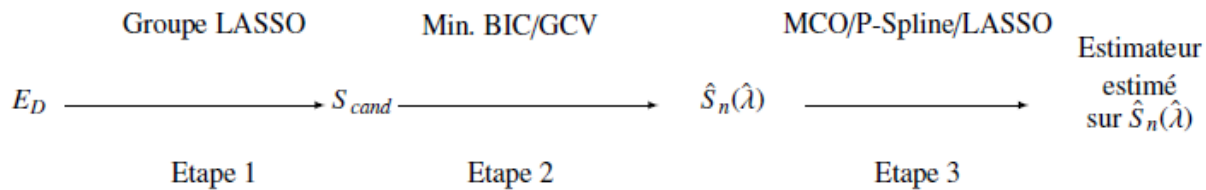


FIGURE 1.5 – Procédure Ante1, Ante2 et Ante2-bis

1.3.2 Procédures en deux ou trois étapes pour sélectionner, régulariser et estimer les modèles

Nous utilisons deux familles de procédures combinant le Group LASSO et les P-Splines.

1.3.2.1 Procédures Ante

Algorithme Procédure Ante1

1. Utilisation du Group Lasso pour sélectionner les variables. Choix du terme de lissage avec un critère BIC, AIC ou GCV (ce choix n'est pas remis en cause dans la deuxième phase)
 2. Etape de régularisation et d'estimation : utilisation de MCO pour estimer les composantes sélectionnées à l'étape 1
-

Algorithme Procédure Ante2

1. Utilisation du Group Lasso pour sélectionner les variables. Choix du terme de lissage avec un critère BIC, AIC ou GCV (ce choix n'est pas remis en cause dans la deuxième phase)
 2. Etape de régularisation et d'estimation : utilisation de P-Splines pour estimer les composantes sélectionnées à l'étape 1
-

Algorithme Procédure Ante2-bis

1. Utilisation du Group Lasso pour sélectionner les variables. Choix du terme de lissage avec un critère BIC, AIC ou GCV (ce choix n'est pas remis en cause dans la deuxième phase)
 2. Etape de régularisation et d'estimation : utilisation d'une pénalité L1 sur tous les coefficients sélectionnés à l'étape 1 simultanément pour réduire les bases des effets conservés
-

Pour les méthodes Ante, le critère AIC n'apparaît pas dans le reste du document parce que ce critère aboutit à des moins bons résultats en termes de prédiction et de sélection de variables sur les simulations et pseudo-simulations. La Figure 1.5 est une illustration des procédures de la famille Ante. E_D est le design initial. Nous calculons les estimateurs Group LASSO associés à une grille de paramètres de lissage, notée Δ_{liss} , de taille 100 par exemple. Après l'étape de Group LASSO, au plus 100 sous-modèles sont candidats. Ils sont notés $S_{cand} = \{\hat{S}_n(\lambda) | \lambda \in \Delta_{liss}\}$. Pour ces procédures, la partie estimation n'influence pas la partie sélection. Pour la seconde famille de procédures, la phase d'estimation a un impact sur la phase de sélection.

Algorithme Procédure Post1

1. Utilisation du Group Lasso pour sélectionner les variables
2. Pour l'ensemble des sous-modèles sélectionnés sur la grille de paramètres de lissage par le Group LASSO, estimation des coefficients à l'aide de MCO
3. Calcul du BIC, de l'AIC ou du GCV avec ces nouveaux estimateurs et conservation de celui qui minimise le critère choisi

Algorithme Procédure Post2

1. Utilisation du Group Lasso pour sélectionner les variables
2. Pour l'ensemble des sous-modèles sélectionnés sur la grille de paramètres de lissage par le Group LASSO, estimation des coefficients à l'aide de P-Splines
3. Calcul du BIC, de l'AIC ou du GCV avec ces nouveaux estimateurs et conservation de celui qui minimise le critère choisi

1.3.2.2 Procédures Post

En conservant les notations de la Figure 1.5, la Figure 1.6 illustre les procédures de la famille Post. Les procédures Post sont différentes de celles Ante, car elles sélectionnent le sous-modèle qui minimise un critère type BIC, AIC ou GCV après une phase d'estimation avec ou sans régularisation, servant à débaiser le Group LASSO (procédures qui peuvent être caractérisées de Post pour le choix du terme de lissage du Group LASSO). Elles demandent la construction de plus de modèles que les trois premières et le coût informatique risque donc d'être plus élevé. La Figure 1.7 met en parallèle les deux familles de procédures.

Dans le paragraphe suivant, nous illustrons les deux familles sur un exemple.

1.3.2.3 Exemple de procédures Ante2Bic et Post2Bic

Nous notons Ante2Bic et Post2Bic les procédures dans les sous-familles respectivement Ante2 et Post2, les deux utilisant le BIC comme critère de sélection de modèles. Dans l'exemple suivant, la charge consommée peut être expliquée par cinq covariables candidates : la température, la nébulosité, l'instant de la journée, le moment dans l'année et la vitesse du vent. Nous utilisons un modèle de la forme du modèle (1.9). Utiliser les cinq covariables ne nous assure pas une prévision optimale (principe de parcimonie).

La Figure 1.8 présente les différentes étapes des procédures Ante et Post appliquées à l'exemple. La première étape est la même pour les procédures Ante et Post. Pour cette première étape, nous initialisons une grille de paramètres de lissage $\Delta_{\text{lissage}} = \{0 < \lambda_{\min} < \dots < \lambda_{\max} < \infty\}$ et calculons les estimateurs Group LASSO associés à chaque élément de Δ_{lissage} . Une variable est sélectionnée lorsque la norme 2 des coefficients estimés par Group LASSO qui lui sont associés est différente de 0. Chaque $\lambda_i \in \Delta_{\text{lissage}}$ conduit à un sous-design candidat, qui peut être redondant

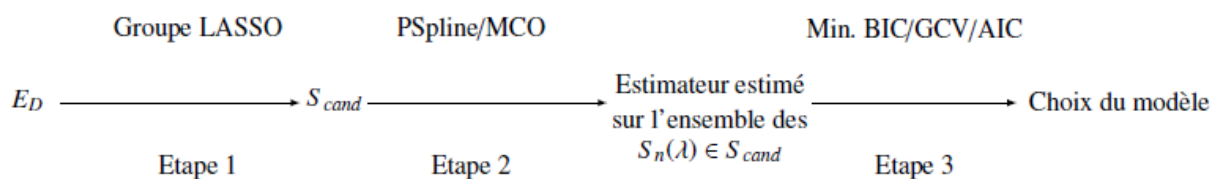


FIGURE 1.6 – Procédure Post1 et Post2

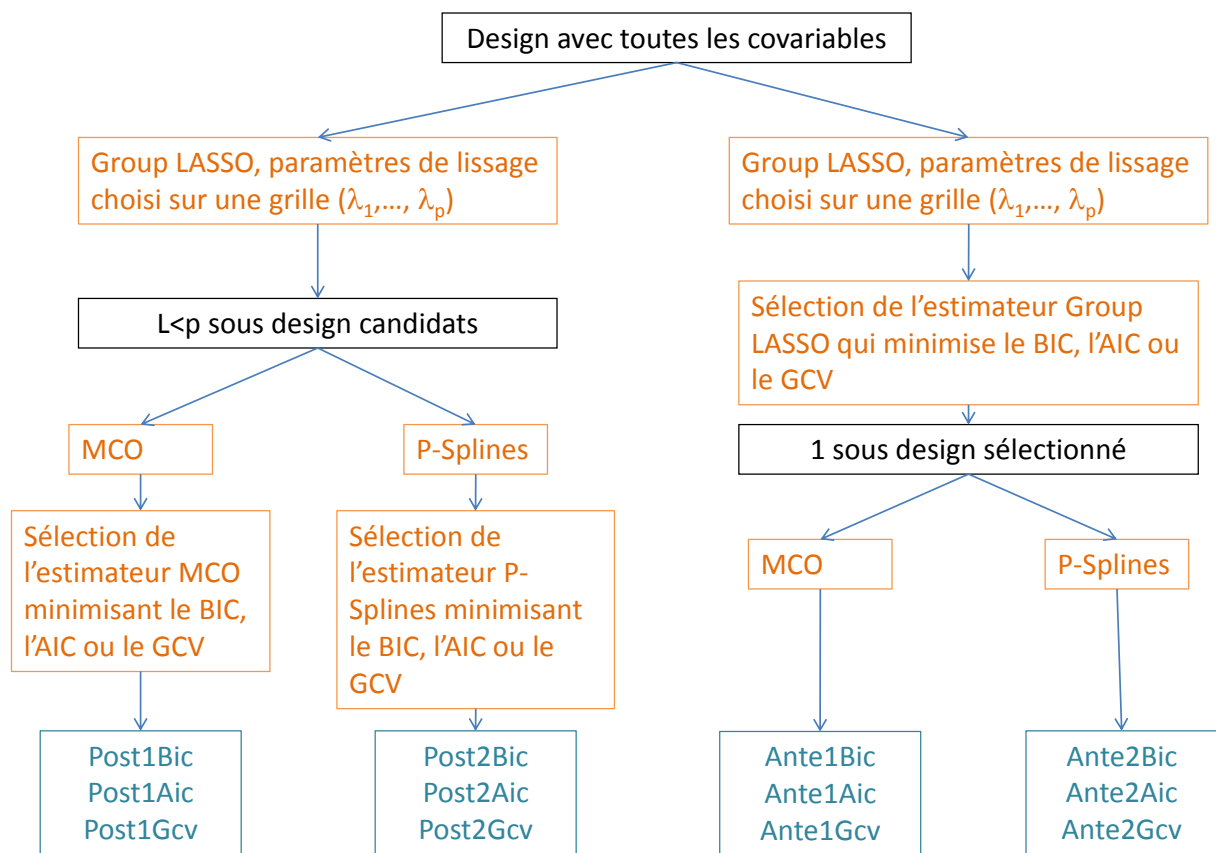


FIGURE 1.7 – Illustration des procédures type Post (gauche) et type Ante (droite)

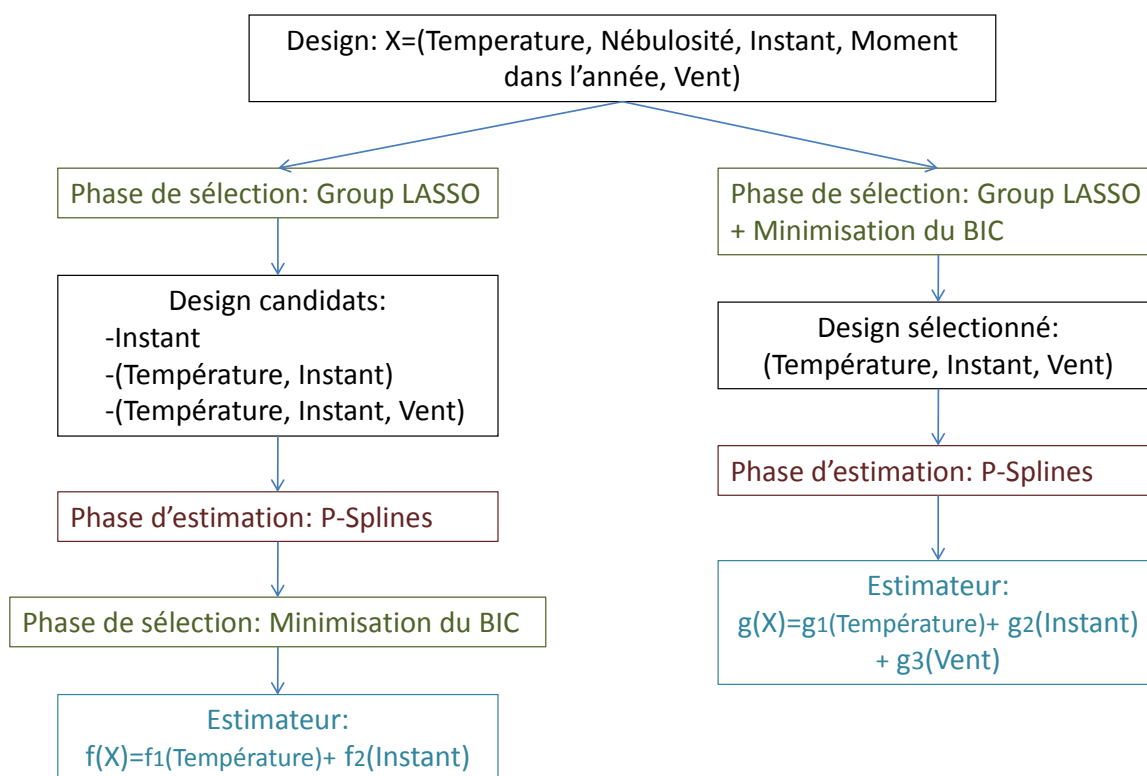


FIGURE 1.8 – Exemple de procédures Post2Bic (gauche) et Ante2Bic (droite)

avec celui obtenu par un autre élément de Δ_{lissage} . Sur notre exemple, trois sous-design candidats sont sélectionnés à partir de l'étape de Group LASSO.

La différence entre Ante2Bic et Post2Bic réside dans le choix du sous-design à utiliser pour obtenir l'estimateur final. Pour la première procédure, le sous-design associé à l'estimateur Group LASSO minimisant le BIC est conservé. Ici, le sous-design sélectionné a trois covariables : (*Temperature, Instant, Vent*). Ensuite, nous estimons le sous-modèle associé à l'aide de l'estimateur P-Splines.

Pour les procédures Post, tous les sous-design candidats que le Group LASSO a permis de construire sont estimés avec des estimateurs P-Splines. L'estimateur qui minimise le BIC et le sous-modèle associé sont sélectionnés. Ici, le design est composé des covariables (*Temperature, Instant*).

1.3.2.4 Nom des méthodes

Le préfixe “Ante” est accolé au nom des méthodes issues de l'une des deux premières procédures, et “Post” pour celles issues des deux dernières. Le degré de liberté est estimé en comptant le nombre de coefficients non nuls.

Pour les simulations, nous nous intéressons aussi au Group LASSO seul et les méthodes de références sont les MCO et les P-Splines utilisant l'ensemble des variables (influentes ou non) comme design. Nous fournissons les résultats que nous obtenons lorsque nous utilisons la fonction *gam* du package **mgcv** avec les options *select=T* (*GAMSelect*) et avec “cs” comme lisseur (*GAMShrinkage*).

La Table 1.3 donne le dictionnaire des méthodes utilisées. Nous donnons deux exemples pour en illustrer la lecture.

Lorsque Ante2Bic est utilisé, *BIC(Group LASSO)+P-Splines* signifie que le plan d'expérience est sélectionné après l'étape Group LASSO grâce au BIC, le modèle étant ensuite ré-estimé à l'aide de P-Splines (méthode de type Ante).

Lorsque Post2Bic est utilisé, *BIC(Group LASSO+P-Splines)* signifie que tous les plans d'expérience issus du Group LASSO sont estimés à l'aide de P-Splines et l'estimateur qui minimise le BIC est sélectionné (méthode de type Post).

Modèle	Méthode d'estimation
Ante1Bic	BIC(Group LASSO)+LASSO
Ante1Gcv	GCV(Group LASSO)+LASSO
Ante2Bic	BIC(Group LASSO)+P-Splines
Ante2Gcv	GCV(Group LASSO)+P-Splines
Post1Bic	BIC(Group LASSO+MCO)
Post1Aic	AIC(Group LASSO+MCO)
Post1Gcv	GCV(Group LASSO+MCO)
Post2Bic	BIC(Group LASSO+P-Splines)
Post2Aic	AIC(Group LASSO+P-Splines)
Post2Gcv	GCV(Group LASSO+P-Splines)
GrpLASSOBic	BIC(Group LASSO)
GrpLASSOGcv	GCV(Group LASSO)
MCO	MCO
PSP	P-Splines
GAMSelect	P-Splines modifiées 1
GAMShrinkage	P-Splines modifiées 2

TABLE 1.3 – Dictionnaire des méthodes utilisées

Dans la section suivante, nous testons nos procédures sur des simulations et des pseudo-simulations.

1.4 Simulations

L'objectif des simulations est d'étudier les trois propriétés suivantes de nos procédures :

- la capacité à sélectionner les composantes ;
- la capacité à reproduire la forme des effets ;
- la capacité à prévoir.

Les questions subsidiaires à ces trois propriétés sont :

- L'intensité du bruit influence-t-elle la qualité de sélection et d'estimation des modèles ? Si oui, comment ?
- Comment le nombre d'observations influence-t-il la qualité de sélection et d'estimation des modèles ?
- Quel impact a la corrélation entre les variables explicatives ?

Nous commençons par présenter les critères d'évaluation des performances.

1.4.1 Critères d'évaluation des performances

Les critères d'évaluation des modèles usuels sont présentés dans la Table 1.4 :

Critère	Définition
MAPE	Mean Absolute Percentage Error ; $MAPE(f, \hat{f}) = 1/n \sum_{i=1}^n (f(x_i) - \hat{f}(x_i))/f(x_i) $
RMSE	Root Mean Square Error ; $RMSE(f, \hat{f}) = \sqrt{1/n \sum_{i=1}^n ((f(x_i) - \hat{f}(x_i))^2)}$
MS	Médiane du nombre de composantes sélectionnées
MTZ	Médiane du nombre de composantes nulles parmi les vraies composantes nulles (vrais zéros)
MFZ	Médiane du nombre de composantes nulles parmi les vraies composantes non nulles (faux zéros)
MTP	Médiane du nombre de composantes non nulles parmi les vraies composantes non nulles (vrais positifs)
MFP	Médiane du nombre de composantes non nulles parmi les vraies composantes nulles (faux négatifs)
Median Curve 1	Courbe estimée correspondant à la simulation associée au RMSE médian sur toutes les simulations pour une méthode donnée

TABLE 1.4 – Critères d'évaluation des performances

MTZ permet de juger les méthodes sur leur capacité à bien seuiller à 0 les composantes nulles. MTP permet d'évaluer leur capacité à ne pas seuiller à 0 des composantes non nulles. MFP permet de quantifier le nombre de composantes que la méthode a conservé alors qu'elles étaient non influentes. Enfin, MFZ permet de quantifier le nombre de composantes que la méthode a seuillé à 0 alors qu'elles étaient influentes. La Table 1.5 résume leur signification.

	Composantes réellement influentes	Composantes réellement nulles
Composantes sélectionnées	MTP	MFP
Composantes non sélectionnées	MFZ	MTZ

TABLE 1.5 – Critères d'évaluation des performances en sélection

1.4.2 Signal to Noise Ratio (SNR)

Le rapport signal sur bruit (SNR) mesure l'intensité du signal par rapport au bruit. Nous utilisons un SNR empirique, qui dépend donc du plan d'expérience utilisé. Nous considérons le modèle suivant :

$$Y = \sum f_j(X_j) + \varepsilon$$

où ε est un bruit centré de variance σ^2 . Soit x_{ji} la i ème observation de la variable X_j . Le SNR est défini par l'équation (1.17) :

$$SNR = \sqrt{\frac{1/(n-1) \sum_{j=1}^d \sum_{i=1}^n (f_j(x_{ji}) - 1/n \sum_{i=1}^n f_j(x_{ji}))^2}{\sigma^2}}. \quad (1.17)$$

Le SNR permet de comparer le signal non bruité avec le bruit. Plus le SNR est fort, moins le signal est bruité, et donc plus il est facile à estimer. Nous considérons que le signal est fortement bruité lorsqu'il est entre 2 et 3 et faiblement bruité s'il est supérieur à 4.

1.4.3 Simulation 1

Description de la simulation Cette simulation s'inspire de la simulation de Lin et Zhang [83] et a aussi été utilisée par Cantoni *et al.* [25] et Antoniadis *et al.* [9]. Parmi les dix covariables candidates, quatre sont informatives. Les composantes de la simulation sont les suivantes :

$$\begin{aligned} f_1 : x &\mapsto 5x \\ f_2 : x &\mapsto 3(2x - 1)^2 \\ f_3 : x &\mapsto \frac{4 \sin(2\pi x)}{2 - \sin(2\pi x)} \\ f_4 : x &\mapsto 6(0.1 \sin(2\pi x) + 0.2 \cos(2\pi x) + 0.3 \sin^2(2\pi x) + 0.4 \cos^3(2\pi x) + 0.5 \sin^4(2\pi x)) \\ f_5, f_6, f_7, f_8, f_9, f_{10} : x &\mapsto 0 \end{aligned}$$

Les cinq fonctions sont $C^\infty[0, 1]$.

La corrélation entre les covariables peut varier. La simulation des covariables se fait selon le principe suivant :

1. Générer $W_{1i}, \dots, W_{10i}, U_i$ indépendamment d'une loi uniforme sur $[0, 1]$, avec $i \in \llbracket 1, n \rrbracket$
2. $X_{ji} = \frac{W_{ji} + tU_i}{1+t}$ avec $j \in \llbracket 1, 10 \rrbracket$

t contrôle la corrélation entre les variables. En effet, $Cor(X_{ji}, X_{li}) = \frac{t^2}{1+t^2}$ pour $j \neq l$.

Le modèle simulé est :

$$Y_i = f_1(X_{1i}) + f_2(X_{2i}) + f_3(X_{3i}) + f_4(X_{4i}) + \sum_{j=5}^{10} f_j(X_{ji}) + \varepsilon_i,$$

où ε_i est un bruit gaussien centré.

Les variables explicatives sont projetées dans des bases de B-Splines de degré 3 et qui ont 10 noeuds placés au niveau des déciles des variables explicatives.

Nous simulons neuf variantes de la simulation avec $n = 3000$ observations. Nous prenons un nombre élevé d'observations parce que les applications concrètes étudiées en contiennent beaucoup. Nous considérons trois valeurs de variance du bruit (0.5, 1.5 et 3.5) et trois valeurs de corrélation ($t = 0, 1$ et 2 , conduisant à une corrélation entre les covariables de respectivement $0, 1/2$ et $4/5$). Pour

chaque fonction, $SNRf_j = \sqrt{\frac{1/(n-1) \sum_{i=1}^n (f_j(x_{ji}) - 1/n \sum_{i=1}^n f_j(x_{ji}))^2}{\sigma^2}}$ quantifie l'influence de la composante considérée.

Simulation	SNR	$SNRf_1$	$SNRf_2$	$SNRf_3$	$SNRf_4$
Variance 0.5, t=0	5.29	2.04	1.25	2.59	3.94
Variance 1.5, t=0	3.09	1.17	0.73	1.49	2.32
Variance 3.5, t=0	2.02	0.78	0.48	0.97	1.51
Variance 0.5, t=1	5.42	1.44	0.85	2.55	4.49
Variance 1.5, t=1	3.12	0.84	0.48	1.49	2.56
Variance 3.5, t=1	2.05	0.55	0.32	0.97	1.69
Variance 0.5, t=2	5.38	1.54	0.83	2.67	4.33
Variance 1.5, t=2	3.10	0.88	0.48	1.53	2.50
Variance 3.5, t=2	2.04	0.58	0.32	1.01	1.64

TABLE 1.6 – SNR des différentes simulations

La Table 1.6 montre que les quatre composantes n'ont pas la même influence sur la variable à expliquer. Les composantes 4 et 2 sont respectivement la plus et la moins influentes. Nous introduisons des variables ayant des influences d'intensité différente pour vérifier que les méthodes parviennent à sélectionner toutes les composantes, même faiblement influentes. Ajouter la corrélation entre les variables explicatives change leur distribution et ainsi les $SNRf_j$ et les SNR .

Les deux modifications des P-Splines proposées par Wood et Marra [86] ont été testées. Elles permettent de réduire les coefficients associés à certaines variables non informatives, mais aucun estimateur n'est strictement nul. Nous avons alors décidé qu'une composante est nulle si la norme 2 de ces coefficients associés est inférieure à un certain seuil. Sur les simulations, les normes 2 sont soit extrêmement faibles (de l'ordre de 10^{-15}), soit proches de 10^{-2} .

Avant de présenter les résultats obtenus par chaque méthode, nous faisons un retour rapide sur l'utilisation des tests pour sélectionner les composantes.

Illustration de l'utilisation des tests sur cette simulation Dans ce paragraphe, nous appliquons dans une méthode de type backward le test de déviance testant la nullité d'une composante dans un modèle additif proposé par Fan et Jiang [44]. La loi de la statistique de test étant inconnue, elle est approchée par une méthode de type bootstrap.

Dans un premier temps, nous considérons des covariables non corrélées avec variance des bruits de 1.5 et 3000 observations. Nous testons

$$H_0 : f_1(x_1) = 0 \text{ vs } H_1 : f_1(x_1) \neq 0 \text{ et } H_0 : f_{10}(x_{10}) = 0 \text{ vs } H_1 : f_{10}(x_{10}) \neq 0.$$

Nous réalisons aussi une sélection backward en présence de covariables non corrélées (sélection backward 1) et en présence de covariables faiblement corrélées (sélection backward 2). Le niveau de confiance est de 0.95. La Table 1.7 récapitule le nombre d'étape associé au bootstrap utilisé pour les trois expérimentations. Nous donnons aussi le temps de calculs approximatif en minutes, la P-Value (pour les deux premières expériences) et les conclusions. Lorsque l'hypothèse H_0 est rejetée, la nullité de la composante est rejetée avec un niveau de confiance de 0.95. Lorsque l'hypothèse de nullité n'est pas rejetée, nous ne pouvons pas rejeter la nullité de la composante avec un niveau de confiance de 0.95.

	Nombre d'étape bootstrap	Temps de calculs	P-Value	Conclusion
Test nullité f_1	200	5 minutes	0.00	Rejet de H_0
Test nullité f_{10}	200	5 minutes	0.46	Non rejet de H_0
Sélection backward 1	50	51 minutes		0 faux négatif 0 faux positif
Sélection backward 2	50	52 minutes		0 faux négatif 1 faux positif

TABLE 1.7 – Tests de nullité lorsque la variance est égale à 1.5

Le temps de calculs excessif et la présence de faux positifs lorsque les covariables sont corrélées confortent notre choix d'utiliser d'autres méthodes. Dans le paragraphe suivant, nous évaluons les capacités de sélection des méthodes comparées.

Capacité de sélection de composantes La Table 1.8 récapitule les critères d'évaluation des capacités de sélection dans le cas où la variance du bruit est de 1.5 et où $t = 2$. Nous donnons la moyenne des critères pour les 200 simulations ainsi que les écart types entre parathèse. Des tables semblables pour les huit autres scénarios sont données dans l'Annexe B.1. Le seul scénario présenté dans ce chapitre est le plus clivant si nous considérons à la fois les capacités de sélection et de prédiction. Si nous avons choisi un scénario avec une variance plus forte, les clivages auraient été plus forts en termes de prédiction, mais plus faibles en termes de sélection. Post1Aic et Post2AIC n'apparaissent pas dans la Table 1.8 par souci de lisibilité. Elles sont légèrement moins bonnes en termes de sélection que respectivement Post1Gcv et Post2Gcv. De même, GAMSelect et GAMShrinkage n'apparaissent pas dans la table, car les deux procédures sont moins performantes.

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	Idéal
MS	4(0)	4(0.26)	4(0)	4(0.16)	4(0.58)	4(0.77)	4
MTZ	6(0)	6(0.26)	6(0)	6(0.16)	6(0.58)	6(0.77)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.26)	0(0)	0(0.16)	0(0.58)	0(0.77)	0

TABLE 1.8 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 1.5 lorsque $t = 2$ (SNR=3.1)

Nous commentons ici les résultats de la Table 1.8 et des tables fournies en Annexe B.1. Aucune variante de la simulation ne comporte pas de faux négatifs. Le vrai modèle est à chaque fois sélectionné par les méthodes Post utilisant le BIC, ce qui n'est pas le cas avec les autres procédures. La méthode GrpLASSOGcv, et donc les méthodes Ante utilisant le GCV, est la plus sensible aux variations du SNR (à corrélation constante, plus celui-ci est faible, plus la méthode est proche de sélectionner à chaque fois le vrai modèle). Les méthodes Ante sont plus impactées par les variations de la corrélation entre les covariables que les méthodes Post. Les tables B.1 à B.8 de l'Annexe B.1 montrent que les méthodes Post utilisant le BIC sont plus performantes en termes de sélection que les méthodes Ante.

Puisque les composantes influentes sont toujours sélectionnées, le nombre de composantes sélectionnées est représentatif des performances de sélection. La Figure 1.9 fournit les nombres de composantes pour les 200 simulations des différentes méthodes lorsque la variance est de 1.5 et $t = 2$. Sur ces figures, les méthodes sont classées selon le nombre de fois où quatre composantes ont

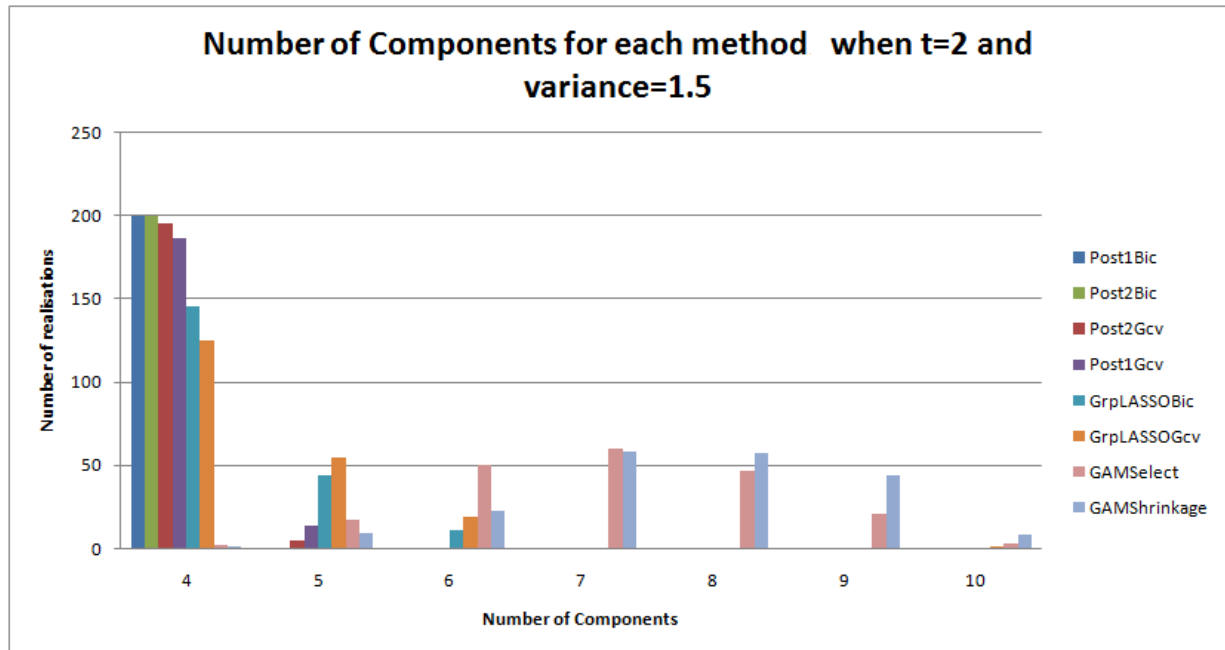


FIGURE 1.9 – Nombre de composantes lorsque la variance du bruit est 1.5 et $t = 2$ (SNR=3.1) pour chaque méthode

été sélectionnées, du plus grand au plus petit (et donc du plus au moins efficace). GAMSelect et GAMShrinkage sont les moins régulières. Post1Bic et Post2Bic ont sélectionné le nombre idoine de composantes pour les 200 simulations. La phase de ré-estimation avant la sélection améliore celle-ci. Le GroupLASSOGcv a sélectionné une fois les dix covariables.

Capacité d'estimation et de prédiction Les formes des courbes des effets sont relativement bien reproduites : par exemple, une parabole est estimée par une parabole. Cependant, l'échelle est moins bien respectée lorsque le Group LASSO est utilisé seul, à cause du biais introduit par l'estimateur. Les étapes suivantes sont donc justifiées. Pour s'affranchir des effets de bord, nous choisissons de représenter sur la Figure 1.10 les courbes estimées provenant de la simulation ayant conduit aux RMSE médians pour GrpLASSOBic (courbe rose), Ante2Bic (courbe noire) et Post2Bic (courbe bleue) lorsque $t = 2$ et pour une variance des bruits de 1.5, en restreignant arbitrairement la variable X_2 à l'intervalle $[0.2, 0.8]$. La courbe rouge représente la vraie courbe. Les effets estimés ne sont pas entièrement lisses. La forme de la courbe pour la composante 2 issue de Post2Bic est plus lisse que les autres courbes estimées. La variance est plus élevée pour GroupLASSOBic et Ante2Bic.

La Figure 1.11 récapitule les boîtes à moustaches des RMSE en prédiction lorsque la variance du bruit est 1.5 et $t = 2$. Les mêmes figures pour les huit autres scénarios sont fournies dans l'Annexe B.2. Dans cet annexe, une figure fournit les boîtes à moustaches rassemblant toutes les méthodes pour le cas où la variance est 1.5 et $t = 2$. Nous n'avons représenté qu'une méthode de référence, qui est généralement le meilleur des modèles de références, c'est-à-dire PSP. Nous avons quatre familles de procédures (Post1, Post2, Ante1 et Ante2). Pour chacune de ces procédures, nous gardons celle utilisant le BIC, afin d'avoir un représentant unique par famille et parce que le BIC est dans la majorité des cas le critère conduisant aux meilleures performances en termes de sélection de composantes à l'intérieur d'une famille de méthodes.

La Figure 1.11 illustre l'importance de la sélection de variables pour la phase de prédiction. En effet, Post2Bic, qui utilise la même méthode d'estimation qu'Ante2Bic et PSP, mais qui est plus performante en termes de sélection, obtient des meilleures performances en termes de prédiction.

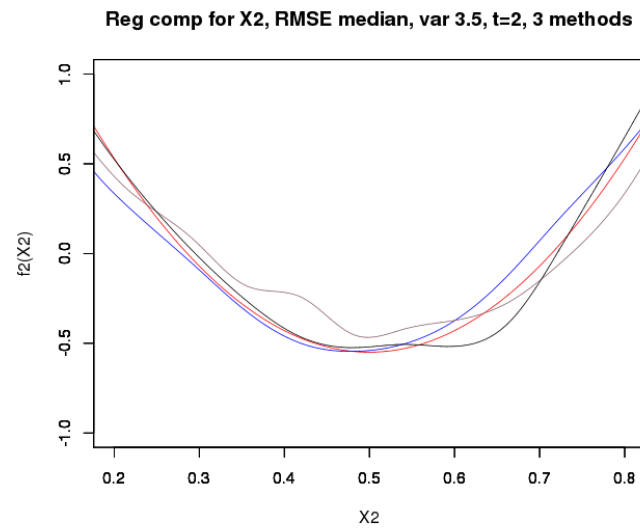


FIGURE 1.10 – Zoom sur les courbes des effets estimés de X_2 lorsque la variance du bruit vaut 1.5 et $t = 2$ (SNR=3.1) pour les modèles amenant aux RMSE médians pour Post2Bic (bleue), GrpLASSOBic (rose) et Ante2Bic (noire)

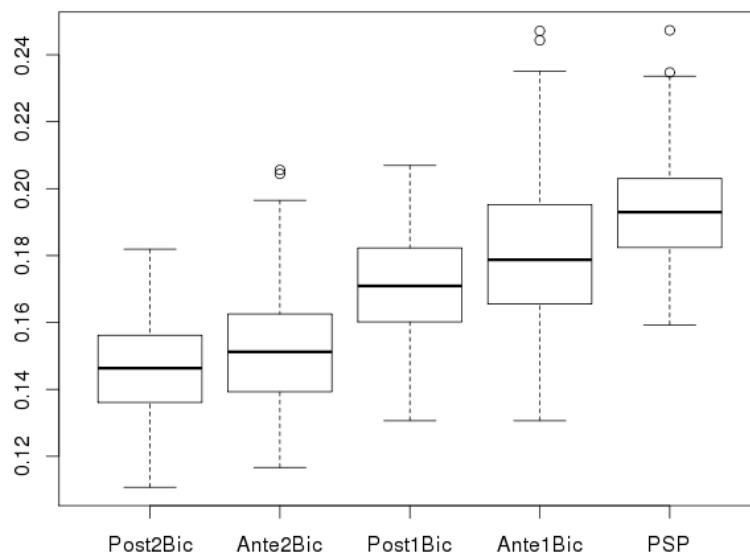


FIGURE 1.11 – Boîtes à moustaches des RMSE lorsque la variance du bruit est 1.5 et $t = 2$ (SNR=3.1)

De même, Post1Bic est plus performante en termes de prédiction qu'Ante1Bic. Sur ce jeu de simulation, la régularisation est importante, car Post2Bic et Ante2Bic sont plus performantes que Post1Bic et Ante1Bic.

Lorsque les variables ne sont pas corrélées, quelle que soit la variance, lorsque la médiane des RMSE est considérée, quatre groupes de méthodes ayant des médianes de RMSE proches se dégagent. Par la suite, nous décrivons ces groupes par performances croissantes. Le premier groupe est composé des méthodes ayant une phase de régularisation avec de la sélection de variables, le deuxième des méthodes sans régularisation avec sélection de variables, le troisième de la méthode de régularisation sans sélection de variables et le quatrième des MCO sans sélection de variables et des méthodes utilisant du Group LASSO sans étape de correction de biais. Enfin, en cas de variance faible des bruits, les méthodes GAMSelect et GAMShrinkage ont une médiane des RMSE forte (à un niveau équivalent du MCO sans sélection de variables). Par contre, lorsque la variance des bruits est forte, les méthodes GAMSelect et GAMShrinkage sont moins dégradées en prédiction que les autres procédures, tout en restant moins efficaces que les méthodes utilisant de la régularisation avec sélection de variables. Les RMSE augmentent lorsque la variance des bruits augmente. Les méthodes avec régularisation et avec sélection de variables sont les plus performantes en termes de prédiction.

Lorsque les covariables sont corrélées, les mêmes quatre groupes se dégagent en termes de RMSE. La méthode Post2Bic devient à chaque fois meilleure que la méthode Ante2Bic, excepté lorsque la variance des bruits et la corrélation entre les covariables sont faibles. Par rapport aux autres méthodes, les performances de GAMSelect et GAMShrinkage sont plus dégradées lorsque les covariables sont corrélées. L'utilisation du GCV peut échouer en termes de RMSE et Ante2Gcv est moins performants en termes de RMSE que Ante1Bic et Post1Bic lorsque $t = 2$ et que la variance des bruits est de 0.5. Avec une corrélation des covariables plus forte ou une variance des bruits plus élevée, les différences entre les méthodes augmentent. Pour les neuf variantes de la simulation, la méthode Post2Bic conduit dans sept cas à la médiane des RMSE la plus faible parmi toutes les méthodes étudiées.

Conclusion Sur les simulations présentées, le choix du critère de sélection influence peu la capacité de sélection des méthodes Post. Pour les méthodes Ante, le GCV est plus dépendant du bruit ou de la corrélation des covariables que le BIC. Dans le cas de variables explicatives fortement corrélées ou de variance du bruit forte (SNR faible), les méthodes Post ont des meilleures performances en termes de sélection que les méthodes Ante. Plus les conditions sont favorables (faible corrélation, SNR fort), plus les méthodes Ante s'approchent des capacités de sélection des méthodes Post. Les méthodes reproduisent fidèlement les familles des formes des effets estimés, mais dans le cas de fortes corrélations entre les covariables, une bonne sélection des composantes est nécessaire pour améliorer l'estimation des effets. La sélection et de l'utilisation de la régularisation améliorent la capacité de prédiction. Dans le cas de conditions défavorables (fortes corrélations entre les variables explicatives, SNR faible), les méthodes Post2 profitent de leur meilleure capacité en termes de sélection par rapport aux méthodes Ante2.

1.4.4 Pseudo-simulation sur des données proches d'EDF

Description de la simulation Nous avons utilisé les données de consommation électrique française pour construire un simulateur de courbes de charge s'inspirant du modèle Métémore [21] (modèle de prévision de consommation d'électricité historique d'EDF). Le modèle de la simulation s'écrit de la manière suivante :

$$\mathbf{Y} = \mathbf{f}_1(\mathbf{X}_1) + \mathbf{f}_2(\mathbf{X}_2) + \mathbf{f}_3(\mathbf{X}_3) + \sum_{i=4}^{14} \mathbf{f}_i(\mathbf{X}_i) + \boldsymbol{\varepsilon},$$

où

- $\mathbf{X}_1 \in \llbracket 0, 23 \rrbracket$ selon l'heure de la journée
- $\mathbf{X}_2 \in \llbracket 1, 8760 \rrbracket$ moment dans l'année
- $\mathbf{X}_3$ issue d'un historique de températures d'une station météorologique
- $\mathbf{X}_4, \dots, \mathbf{X}_{13}$ issues de dix historiques de températures de dix stations météorologiques différentes (et différentes de la station météorologique dont est issue $\mathbf{X}_3$)
- $\mathbf{X}_{14}$ issue d'un historique de vitesses du vent
- ε bruit gaussien centré

Ici, les variables sont très fortement corrélées, avec une corrélation maximale de 0.98. Il y a un risque de concurvité qui est l'équivalent pour les modèles additifs des problèmes de multicollinéarité pour les modèles linéaires et qui conduit à des problèmes d'identifiabilité ainsi que des biais dans l'estimation des paramètres du modèle et de l'erreur. Ce problème peut être réduit en sélectionnant les covariables.

Les applications associées au modèle sont les suivantes :

$$\begin{aligned}
 f_1 : x &\mapsto \sum_{h=1}^3 (a_h \sin(h\omega_1 x) + b_h \cos(h\omega_1 x)), \\
 f_2 : x &\mapsto \sum_{h=1}^4 (c_h \sin(h\omega_2 x) + d_h \cos(h\omega_2 x)), \\
 f_3 : x &\mapsto 380 f_{tempChaufage}(x - 12.5) + 40 f_{tempClim}(x - 20), \\
 f_4 : x &\mapsto 0,
 \end{aligned}$$

où a_h, b_h, c_h et d_h coefficients estimés sur des données EDF, $\omega_1 = \frac{\pi}{23}$, $\omega_2 = \frac{\pi}{8760}$ et

$$\begin{aligned}
 erf : x &\mapsto \frac{2}{\sqrt{\pi}} \int_0^x \exp(-t^2) dt, \\
 f_{tempChaufage} : x &\mapsto \sqrt{\frac{2}{\pi}} \exp(-x^2/8) - \frac{1}{2} x \left(1 - erf\left(\frac{x}{2\sqrt{2}}\right) \right), \\
 f_{tempClim} : x &\mapsto \sqrt{\frac{2}{\pi}} \exp(-x^2/8) + \frac{1}{2} x \left(1 - erf\left(\frac{-x}{2\sqrt{2}}\right) \right).
 \end{aligned}$$

f_1 , f_2 et f_3 sont C^∞ .

Nous appelons cette pseudo-simulation “simulation type EDF 1”. La Figure 1.12 donne les courbes des effets des variables *Temperature*, *Instant* et *Toy* du simulateur.

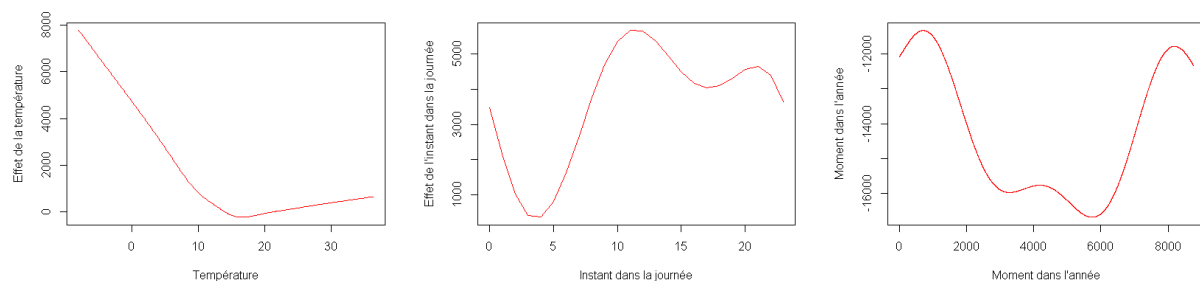


FIGURE 1.12 – De gauche à droite, courbes des effets des variables *Temperature*, *Instant* et *Toy* de la simulation type EDF 1

La Figure 1.13 donne une courbe de charge non bruitée simulée. Elle sert de référence car le SNR varie.

Simulation	SNR	SNR_{f_1}	SNR_{f_2}	SNR_{f_3}
Variance 5.10^5	4.40	2.32	2.69	2.60
Variance 1.10^6	3.11	1.64	1.90	1.84
Variance 2.10^6	2.20	1.16	1.35	1.30

TABLE 1.9 – SNR des différentes pseudo-simulations

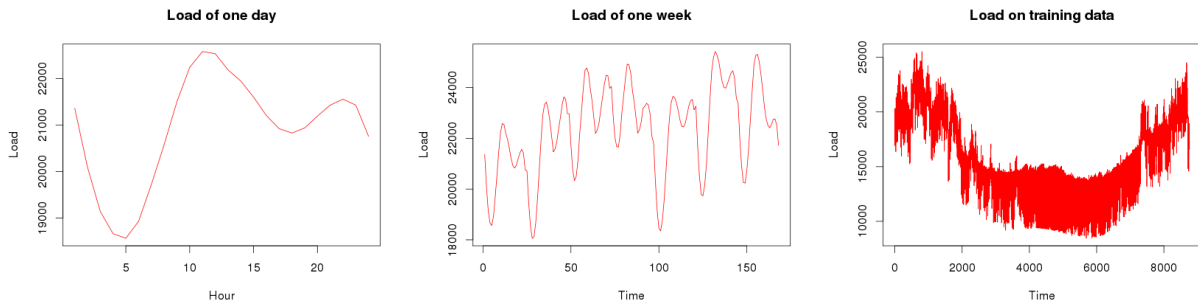


FIGURE 1.13 – De gauche à droite, charge simulée sur un jour, une semaine et une année de la simulation type EDF 1

Nous reconnaissons sur la Figure 1.13 des phénomènes réalistes pour la consommation électrique française. En effet, les trois saisonnalités journalière, hebdomadaire et annuelle sont présentes et réalistes, tout comme l'effet non linéaire de la température.

Les 14 variables explicatives sont projetées dans des bases de B-Splines de degré 3 à 10 noeuds placés au niveau des déciles. Nous étudions des variantes de la simulation ayant pour variance 5.10^5 , 1.10^6 et 2.10^6 (voir Table 1.9). Chaque variante comporte 200 répétitions.

Qualité de sélection de variables La Table 1.10 donne les critères d'évaluation de la capacité de sélection de variables sur les 200 simulations type EDF 1 lorsque le SNR est 3.11. Les tables B.9 et B.10 en Annexe B.3 fournissent ces mêmes critères pour les deux autres SNR.

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	Idéal
MS	3(0)	3(0.30)	3(0)	3(0.22)	5(1.01)	4(0.66)	3
MTZ	11(0)	11(0.30)	11(0)	11(0.22)	9(1.01)	10(0.66)	11
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	3(0)	3(0)	3(0)	3(0)	3(0)	3(0)	3
MFP	0(0)	0(0.30)	0(0)	0(0.22)	2(1.01)	1(0.66)	0

TABLE 1.10 – Capacité à sélectionner les composantes lorsque la variance du bruit est de 1.10^6 (SNR=3.11) pour la simulation type EDF 1

Les covariables candidates sont très fortement corrélées. Même si la sélection des variables influentes est alors plus compliquée, il n'y a pas de faux négatifs. Les méthodes Ante sont moins performantes que les méthodes Post. Pour celles-ci, le BIC est plus performant en termes de sélection. L'AIC et le GCV ont les mêmes performances en termes de sélection (l'AIC n'apparaît pas dans les tables par souci de lisibilité).

Qualité de prédiction La Figure 1.14 présente la courbe des effets estimés par les méthodes Ante2Bic et Post2Bic pour la simulation ayant le RMSE médian pour chaque méthode.

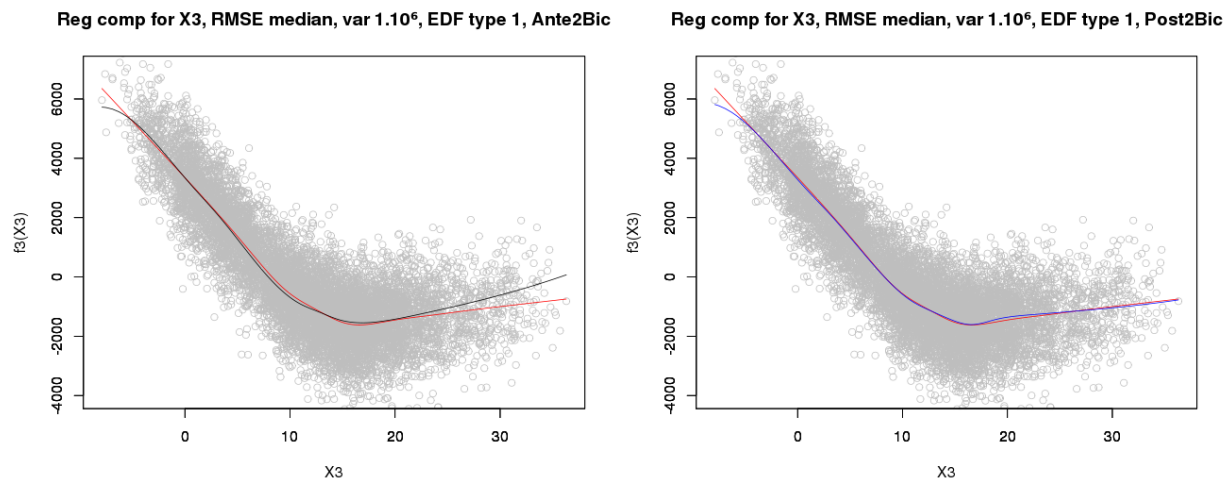


FIGURE 1.14 – Courbes des effets estimés de X_3 lorsque la variance du bruit vaut 1.10^6 (SNR=3.1) pour les modèles conduisant aux RMSE médians

Les méthodes retrouvent la famille de la fonction associée à X_3 mais Ante2Bic reproduit moins fidèlement la courbe pour les températures chaudes. L'explication peut être que cette méthode sélectionne des variables non influentes contrairement à la méthode Post2Bic.

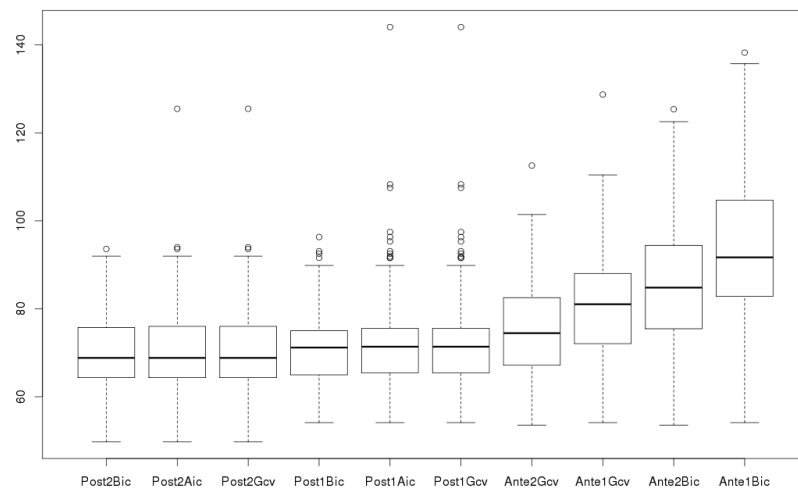


FIGURE 1.15 – Boîtes à moustaches des RMSE lorsque la variance du bruit vaut 1.10^6 (SNR=3.1)

La Figure 1.15 montre que le BIC, qui est moins performant en termes de sélection de variables, est moins performant en termes de prédiction que le GCV pour les méthodes Ante. Les méthodes Post, avec ou sans régularisation, sont plus performantes que les méthodes Ante. Pour les méthodes Post, le BIC conduit aux meilleures prédictions, parce qu'il est plus performant en sélection de covariables.

Bien que les covariables soient très fortement corrélées, les conclusions pour cette pseudo-simulation sont les mêmes que pour la simulation précédente.

1.4.5 Conclusion

Les simulations et les pseudo-simulations présentées comportent beaucoup d'observations. L'absence de faux négatifs conduit à ce que les méthodes soient plus performantes en termes de sélection lorsque le SNR est faible. Le vrai modèle fait partie des modèles candidats. Le BIC est alors connu pour être le critère le plus performant en termes de sélection. Nos simulations et pseudo-simulations illustrent cette propriété. Plus le paramètre de lissage du Group LASSO est fort, moins il y a de variables sélectionnées, et plus son estimateur est biaisé. Ceci explique qu'il y ait des faux positifs pour les méthodes Ante, puisque la sélection de modèles est réalisée dès l'étape de Group LASSO. Il y a donc un compromis entre bonne sélection et biais trop élevé de l'estimateur Group LASSO à trouver pour ces méthodes. Pour les méthodes Post, la sélection est réalisée après que la correction du biais introduit par le Group LASSO, ce qui explique que les méthodes Post soient plus performantes en termes de sélection. Nous avons réalisé les simulations (non présentées dans ce document) dans des cas où la corrélation est extrêmement forte. Au-dessus d'une corrélation entre les covariables de 0.99, les méthodes ne parviennent plus à sélectionner les covariables de manière performante.

L'estimateur obtenu par Group LASSO sans ré-estimation a des moins bonnes performances prédictives que l'estimateur Group LASSO oracle, parce que son paramètre de lissage est plus fort, et donc le biais introduit aussi. Pour les procédures d'estimation en plusieurs étapes, telles que les procédures Ante et Post, l'estimateur obtenu est équivalent à l'estimateur oracle si le vrai design est sélectionné. Post2Bic conduit aux meilleures performances en termes de prédiction par rapport aux autres méthodes.

Les simulations semblent confirmer que les procédures établies sont performantes en termes de sélection et d'estimation. Le travail du chapitre 3 consiste à montrer théoriquement que les procédures Post1Bic rassemblent bien les propriétés de sélection du Group LASSO et d'estimation des MCO. Dans le chapitre 2, nous utilisons des estimateurs en plusieurs étapes sur des données de consommation électrique.

Estimateurs en plusieurs étapes de modèles additifs appliqués à la prévision de consommation à plusieurs niveaux d'agrégation

Introduction

Nous appliquons des méthodes automatiques de sélection de variables pour les modèles additifs servant à prévoir la consommation électrique à différents niveaux d'agrégation. De nombreux articles démontrent l'intérêt de ces modèles pour prévoir la consommation à un niveau national (voir [94]) et local (voir [55]) en France, ainsi que régional en Australie (voir [46]). Trois des dix équipes les mieux classées (voir [17] et [91]) de la compétition GEFCOM 2012 [63] les ont utilisés. Dans la compétition GEFCOM 2014 [64] de prévision probabiliste de la charge électrique consommée, sujet en plein essor du fait du développement des énergies renouvelables et donc de l'importance de la modélisation des aléas, les deux premières équipes [51] et [39] les ont aussi utilisés. Dans les articles cités précédemment, les covariables sont sélectionnées en combinant expertise métier et sélection pas à pas, ce qui est chronophage et peu généralisable. Nous appliquons la méthodologie de sélection automatique de variables pour les modèles additifs présentée dans le chapitre précédent sur plusieurs jeux de données correspondant à des problématiques différentes. Nous proposons une méthodologie utilisant la structure temporelle des données pour améliorer la sélection. Nous abordons également le sujet de la prévision court terme (à un horizon de 24 heures) et adaptons nos méthodes aux modèles linéaires pour sélectionner automatiquement des modèles auto-régressifs. L'article [111] reprend certains résultats du chapitre.

Nous étudions trois jeux de données. Le premier concerne l'ensemble des clients d'EDF (que nous appellerons désormais données agrégées), ce qui est un haut niveau d'agrégation (près de 28 millions de clients en France en 2010). Les deux autres jeux de données concernent des données locales. Elles comportent plusieurs séries temporelles plus ou moins similaires. Ces données contiennent les consommations de différents points d'un réseau électrique. Les données agrégées étant historiquement déjà bien connues, il existe des modèles très performants à EDF pour les prévoir. Plus qu'améliorer les performances, notre objectif est de sélectionner et d'estimer automatiquement des modèles obtenant des performances similaires. L'automatisation est nécessaire pour avoir une plus grande adaptabilité aux changements et pour améliorer la reproductibilité des études. Au niveau local, le nombre important de séries chronologiques (par exemple, pour les données ERDF, environ 2200 postes sources et plus de 750 000 postes de transformation) imposent l'automatisation des méthodes. Ces données sont moins connues que les données agrégées, les rapports de signal sur bruit sont plus faibles et les covariables candidates plus nombreuses.

La section 2.1 présente une application pratique où nous utilisons les procédures automatiques Post2 (voir chapitre 1) sur des données du portefeuille EDF et nous comparons les performances moyen et court terme des modèles obtenus avec celles d'un modèle additif d'un expert EDF R&D

(appelé modèle EDF par la suite), qui sert de fondement à ce travail.

Dans la section 2.2, nous considérons deux jeux de données sur les réseaux de distribution américain (données GEFCom 2012) et français (postes sources). Les données du site GEFCom 2012 sont des données publiques étudiées par des chercheurs et professionnels extérieurs à EDF (voir par exemple [105], [84], [17]), nous permettant ainsi d’avoir des méthodes concurrentes à tester. Ces données représentent un cas d’application intéressant de sélection des localisations des mesures des covariables météorologiques utilisées pour prévoir la consommation locale d’électricité (voir [65]). Les procédures Post2 permettent d’obtenir des modèles moyen terme incorporant un sous-ensemble de stations météorologiques performant en chaque point du réseau. Ensuite, nous adaptons la procédure Post1 (voir chapitre 1) pour sélectionner les erreurs retardées pour la prévision court terme. Une soixantaine de postes sources en France ont aussi étudié avec la procédure Post2.

Pour réaliser le travail présenté dans ce chapitre, nous avons implémenté le prototype de package R **CASA** (Component Automatic Selection in Additive model) permettant d’utiliser les procédures de sélection automatique de composantes dans les modèles additifs présentées dans le chapitre précédent sur des données réelles. Une vignette de ce package, qui a été présenté à la conférence UseR! 2015 à Aalborg, est fournie dans l’Annexe A.

2.1 Modélisation et prévision de la consommation pour le portefeuille EDF

Nous travaillons ici sur le portefeuille EDF. Cet agrégat va connaître des changements majeurs dans les années à venir. En effet, la fin de certains tarifs réglementés en 2016 risque de conduire à de nombreux départs de clients vers la concurrence, notamment dans le secteur des profilés (clients dont la consommation est profilée). De plus, les modèles doivent facilement s’adapter aux changements d’habitudes des consommateurs. Les modèles opérationnels utilisés pour le portefeuille EDF sont majoritairement paramétriques et peu évolutifs, ce qui explique le développement actuel à EDF des modèles additifs, qui répondent bien au besoin d’adaptivité. L’intégration de nouvelles covariables y réalisée facilement. De précédentes études internes à EDF ont montré que ces modèles possèdent notamment une bonne capacité à modéliser la thermosensibilité de la consommation électrique grâce à l’introduction de plusieurs covariables liées à la température réalisée. Ceci motive l’utilisation de nos méthodes automatiques appliquées à la sélection des variables météorologiques pour des modèles additifs. Nous construisons un dictionnaire de covariables météorologiques, issues principalement de la température, mais aussi de la nébulosité et de la vitesse du vent, pour sélectionner des modèles de prévision de la charge consommée sur le portefeuille EDF. L’intégration et l’utilisation du vent (voir [2]) sont des sujets à part entière qui ne sont pas abordés ici.

2.1.1 Description des données

Nous disposons des données des consommations électriques des clients d’EDF entre septembre 2007 et juillet 2013. Nous apprenons les modèles sur la période septembre 2007-août 2012 et les testons entre septembre 2012 et juillet 2013.

Aspects calendaires Dans ce paragraphe, nous illustrons les cycles annuel, hebdomadaire et journalier de la consommation électrique pour le portefeuille EDF. Le graphique de gauche de la Figure 2.1 présente la charge journalière moyenne consommée du portefeuille EDF pendant l’année 2008. Le graphique de droite de la Figure 2.1 représente la charge consommée durant deux semaines en 2008 : une en hiver et l’autre en été.

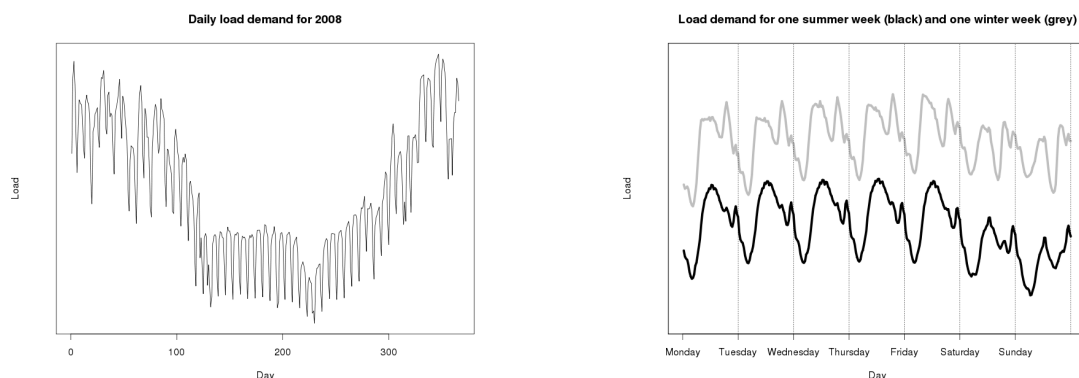


FIGURE 2.1 – Charge du portefeuille consommée en 2008 (gauche) et sur deux semaines (droite) : une pendant l’hiver (grise) et l’autre pendant l’été 2008 (noire)

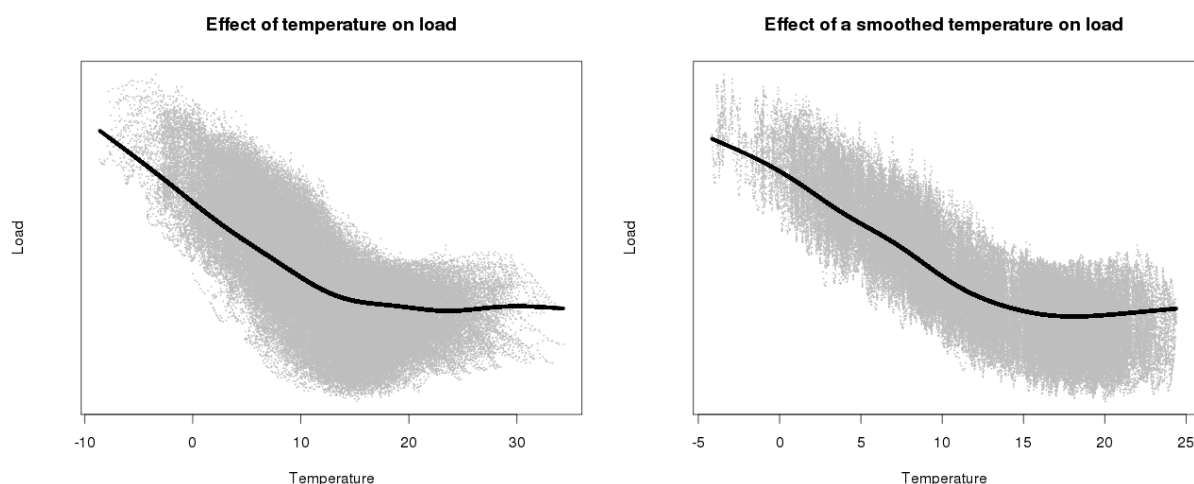


FIGURE 2.2 – Effet estimé de la température et de la température lissée sur la charge consommée en 2008 sur le portefeuille EDF

Le chauffage étant majoritairement électrique en France, la charge consommée est plus élevée en hiver qu’en été. Le cycle annuel s’explique en partie par l’activité économique qui justifie les ruptures de consommation pendant les vacances, notamment autour de Noël et du 15 août. Le niveau de charge consommée est plus fort en hiver qu’en été, illustrant l’impact du climat. Il y a aussi des cycles hebdomadaire (consommation plus forte les jours ouvrables que le week-end) et journalier (consommation plus basse la nuit que le jour). Les pics de consommation ne sont pas identiques en été et en hiver. Celui de 18/19h est plus marqué en hiver. En été comme en hiver, des pics de consommation se réalisent le matin et vers 22h, heure de l’enclenchement des chauffe-eaux.

Le changement d’heure hiver-été influence aussi la consommation électrique. Il a été introduit en 1976 pour réduire la consommation d’électricité, notamment l’été, en jouant sur l’éclairage utilisé en soirée. Il lisse les pics de consommation au printemps et en automne.

Impact climatique La Figure 2.2 présente la charge consommée en fonction de la température brute et d’une température lissée avec un lissage exponentiel simple. Ici, la température correspond à une moyenne pondérée (poids historiques) fournie par Météo France des températures de 32 stations météorologiques françaises (voir Annexe C.1). Le lien entre la température et la charge électrique consommée pour le portefeuille EDF est non linéaire. Sous une certaine température,

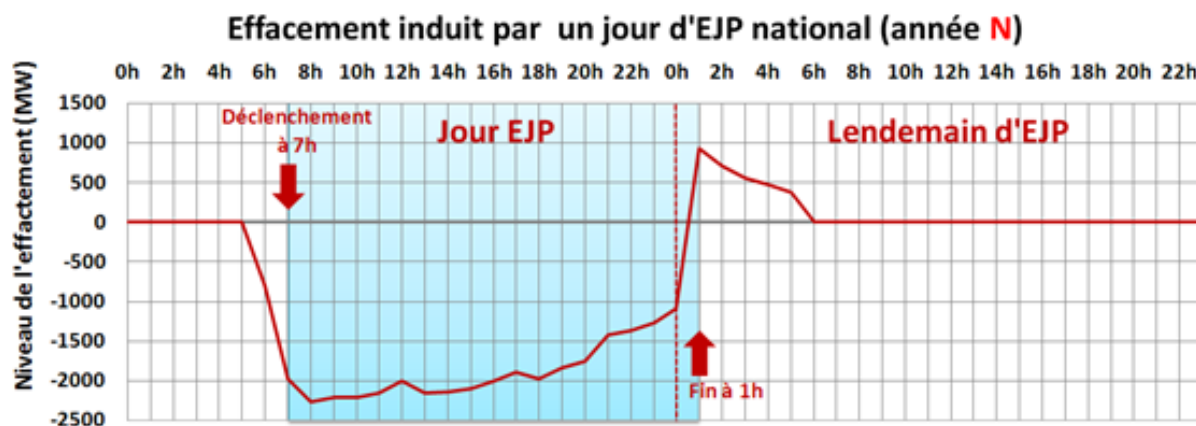


FIGURE 2.3 – Effet d'un effacement (source : Présentation "Les contrats d'effacements tarifaires", séminaire effacement interne EDF 2013, Both/Pichavant)

appelée seuil de chauffage, la consommation électrique augmente lorsque la température baisse, illustrant ainsi l'impact important du chauffage électrique en France. De même, il existe une température, appelée seuil de climatisation, au-dessus de laquelle une hausse de la température conduit à une consommation plus forte. Ce seuil s'explique par l'utilisation de la climatisation. Celle-ci étant moins développée que le chauffage électrique en France, la part climatisation est moins forte que la part chauffage. Ceci n'est pas observé dans certaines régions du monde pour lesquelles la climatisation est plus développée (Floride ou Australie par exemple).

Contrairement aux variables calendaires, les variables météorologiques ne sont pas déterministes. Dans ce document, nous utilisons les données réalisées, même en prévision, ce qui n'est pas réaliste en pratique. Il s'agit néanmoins d'une bonne méthode pour évaluer les performances des modèles conditionnellement à ces covariables. Ces mesures seraient remplacées par des prévisions dans le cas d'une application réelle, rajoutant de l'incertitude. Cependant, dans notre cas, nous nous intéressons principalement à des études comparatives.

Impact tarifaire Il existe des contrats, dits d'effacement tarifaire¹, pour lesquels la tarification de l'électricité est dynamique. Les clients bénéficient alors d'une réduction des prix de l'électricité la majorité des jours, avec en contrepartie un tarif de l'électricité beaucoup plus élevé quelques jours dans l'année, les encourageant ainsi à moins consommer et donc à s'effacer. EDF utilise ces jours lorsque le coût d'approvisionnement en électricité est élevé ou la demande très forte (jour de grand froid par exemple). Le contrat EJP (Effacement Jour de Pointe) fait partie de ces contrats d'effacement. Ce tarif permet de bénéficier pendant 343 jours par an d'une réduction tarifaire proche du tarif des heures creuses. En contrepartie, le prix de l'électricité est nettement plus élevé pendant 22 jours étalés du 1er novembre au 31 mars. Lorsqu'il est déclenché, l'EJP dure 18 heures : il débute à 7h du matin et se termine à 1h le lendemain. L'effet d'un effacement est représenté dans le Figure 2.3. La consommation baisse avant 7h du matin (effet d'anticipation) et il y a un pic de consommation à la fin de l'EJP (effet rebond). Il est compliqué de savoir quand déclencher un effacement, de mesurer ce qu'il peut apporter, de prévoir les niveaux qu'auront les effets d'anticipation, d'effacement et de rebond, ainsi que de quantifier ce qu'il a réellement apporté. Pour EDF, les enjeux financiers et industriels sont importants. Cependant, ce sujet n'entre pas dans le cadre de cette thèse et nous nous affranchissons de ces problèmes de tarification. Les charges ont été préalablement corrigées des effacements. Le lecteur intéressé par ce sujet peut se reporter, entre autres, à la thèse de Leslie Hatton, réalisée à EDF R&D et à l'Université de

1. http://www.developpement-durable.gouv.fr/IMG/pdf/15_-_La_production_d_electricite_et_l_effacement_de_consommation_en_France-Def.pdf

Rennes et soutenue en janvier 2015 (voir [59]).

Saisonnalité journalière : modélisation par instant de la journée Estimer un modèle global signifie que nous estimons un modèle unique commun à tous les pas de temps journaliers de l'étude, alors qu'estimer des modèles selon l'instant (nous appellerons par la suite ces modèles "modèles par instant") signifie que nous estimons un modèle spécifique à chaque unité de pas de temps de la journée. Par exemple, si nous avons des données horaires, nous estimons 24 modèles différents selon l'heure de la journée. Nous utilisons les seconds, car la consommation d'électricité est fortement corrélée avec l'instant de la journée. De plus, les effets changent selon l'instant de la journée : la température n'a pas le même impact sur la consommation d'électricité la nuit et le jour par exemple. La Figure 2.4 illustre ce phénomène.

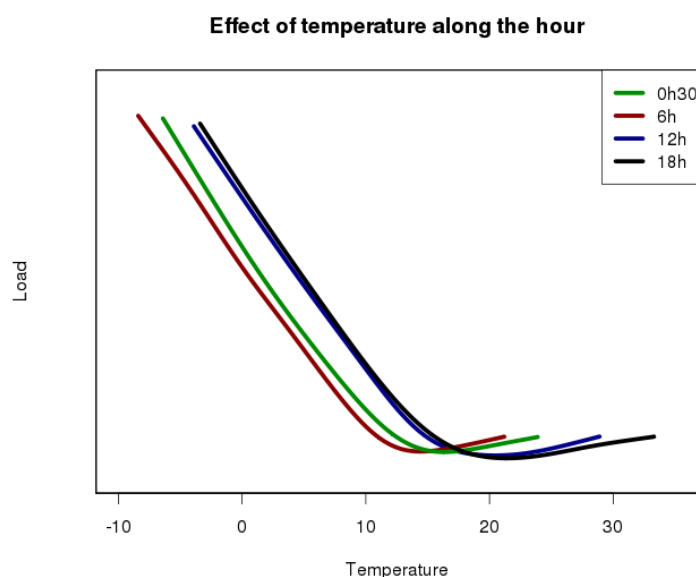


FIGURE 2.4 – Effet de la température selon l'instant de la journée sur la charge consommée

En adoptant un modèle différent selon l'instant de la journée, nous supposons l'indépendance entre chaque instant, ce qui est une hypothèse erronée. Traiter chaque série temporelle indépendamment des autres casse l'aspect longitudinal des données observées. Cependant, en pratique, les experts d'EDF (voir [94] ou [55] par exemple) et hors EDF (voir [46]), tout comme nous, ont constaté que les modèles par instant sont plus performants en prévision. De plus, cette hypothèse semble nécessaire, car si aucune hypothèse sur la structure de la matrice de covariance n'est effectuée, celle-ci serait probablement de grande dimension et donc les modèles additifs compliqués à estimer.

2.1.2 Modèle paramétrique historique d'EDF

L'article de Bruhns *et al.* [21] présente un modèle couramment utilisé par EDF pour la prévision de consommation électrique court et moyen termes. Il se décompose en deux parts :

- une dépendante du climat (principalement de la température) ;
- une indépendante du climat modélisant la tendance et les saisonnalités.

Le modèle s'écrit donc sous la forme suivante :

$$P_i = Pc_i + Phc_i + \varepsilon_i,$$

où

- P_i est la charge consommée pour l'observation i
- Pc_i la partie modélisant la partie climatique de la charge consommée
- Phc_i la partie modélisant la partie hors climatique de la charge consommée
- ε_i bruit supposé gaussien

Pc est estimée non linéairement pour modéliser les effets suivants :

- lorsque la température est plus froide qu'un certain seuil, la charge consommée augmente (gradient de chauffage) ;
- lorsque la température est plus chaude qu'un certain seuil, la charge consommée augmente (gradient de climatisation).

Pc est composée de deux parts qui s'additionnent : la part chauffage, notée $\bar{P}c$, et la part climatisation. Nous présentons seulement la part chauffage dans ce document. Un ensemble S de stations météorologiques est utilisé. Chaque série de températures est lissée exponentiellement. Si on note $t_{i,s}$ la température de l'observation i de la station météorologique s , la température lissée $u_{i,s}$ est telle que $u_{i,s} = \theta u_{i-1,s} + (1 - \theta)t_{i,s}$. Généralement, $\theta \approx 0.98$ lorsque le pas est horaire. Ensuite, cette température lissée est agrégée avec la température observée et la nébulosité, notée $n_{i,s}$:

$$v_{i,s} = (1 - \alpha_h)u_{i,s} + \alpha_h t_{i,s} + \mu_h n_{i,s}.$$

$v_{i,s}$ est l'entrée de la fonction

$$\Psi(v_{i,s}, t, \sigma) = E(\min(v_{i,s} - T, 0)),$$

où t est la température seuil, σ la dispersion autour de ce seuil et $T \sim N(t, \sigma^2)$. Avec cette modélisation, l'effet de la part chauffage est supposé linéaire, ce qui peut être remis en cause pour les températures très faibles.

Finalement, la part chauffage peut s'écrire

$$\bar{P}c_i = \sum_{s \in S} g_{h,s} \times (1 + r \times y_i) \times \Psi(v_{i,s}, t_h, \sigma),$$

où h et y_i dépendent respectivement de l'heure et de la date, r et $g_{h,s}$ sont les gradients respectivement de la tendance et du chauffage pour la station s et l'heure h .

Phc_i se décompose en trois composantes qui sont multipliées entre elles :

- $\pi_{h,k}$, dépendant de l'heure et du type de jours k , modélise les saisonnalités journalières et hebdomadaires ;
- $S_{i,p,h}$, avec p découpant une année en plusieurs sous-périodes de consommation homogène, modélise les saisonnalités annuelles ;
- une tendance $R_i = 1 + r \times y_i$.

S_i se décompose en la somme de plusieurs types de variables : des variables indicatrices, tenant compte des jours fériés et des vacances, et d'une variable qui est la projection dans une base de série de Fourier, tronquée aux quatre premiers termes, du nombre de jours de l'année divisé par 365.25. La décomposition varie selon l'heure. Les types de jours sont obtenus grâce des méthodes de clustering type k-means.

Ce modèle est peu évolutif. Pour ajouter l'effet d'une nouvelle covariable, des méthodes non-paramétriques sont utilisées pour estimer sa forme, puis celle-ci est "reproduite" à l'aide de méthodes paramétriques. Le modèle comporte beaucoup de paramètres (environ 2000), car il estime paramétriquement des effets non linéaires (voir par exemple les Figures 2.1 ou 2.2). Ce nombre de paramètres à estimer nous conduit à utiliser d'autres modèles pour estimer et prévoir la charge consommée. Potentiellement, nous pouvons avoir un grand nombre de covariables, ce qui explique que nous n'utilisons pas de modèle ni de méthode purement non-paramétrique (fléau de la dimension). Le modèle additif permet de réaliser un bon compromis entre la complexité et flexibilité du modèle.

2.1.3 Modèle additif moyen terme d'EDF construit grâce à l'expertise métier

Dans cette sous-section, nous expliquons de manière simplifiée la méthode utilisée à EDF pour construire le modèle auquel nous nous comparons par la suite. Nous commençons par décrire les covariables candidates utilisées pour modéliser la charge consommée.

2.1.3.1 Covariables candidates

Un dictionnaire de covariables candidates servant à modéliser les parts climatique et saisonnière a été construit. Nous le décrivons brièvement.

Variables calendaires Les variables calendaires permettent de modéliser les saisonnalités annuelles ou hebdomadaires. Par exemple, le moment dans l'année (allant de 0 à 1 du 1er janvier au 31 décembre) et le type de jours servent à modéliser respectivement les saisonnalités annuelle et hebdomadaire (niveau de charge différent entre un mardi et un dimanche par exemple). Ces variables déterministes sont classiquement utilisées dans les modèles de prévision de consommation électrique.

Variables météorologiques Nous utilisons les variables météorologiques brutes (température, vitesse du vent et nébulosité instantanées ou retardées de quelques heures ou de quelques jours), agrégées (température moyenne du jour ou de la veille par exemple) et lissées (lissages exponentiels simples). Les versions modifiées de variables météorologiques brutes servent à réduire la variabilité de ces covariables et à modéliser l'effet de l'inertie des bâtiments. À partir de la température, de la vitesse du vent et de la nébulosité, l'expert construit 76 covariables.

Finalement, l'expert utilise ce dictionnaire de plus de 80 covariables pour sélectionner un modèle additif.

2.1.3.2 Sélection de variables par expertise métier

Pour sélectionner un modèle, l'expert utilise des méthodes s'approchant des méthodes de recherche pas à pas : il teste manuellement l'impact sur le critère de validation croisée généralisée (GCV) de l'introduction ou du retrait d'une covariable et teste des choix de facteurs de lissage différents.

L'expert obtient le modèle (2.1). Par souci de lisibilité, nous ne faisons pas apparaître la dépendance en l'instant de la journée dans ce modèle.

$$\begin{aligned}
 Y_t = & \beta_0 + h(t) + f(\text{Theating}_t) \\
 & + \sum_{i=1}^7 \sum_{j=1}^3 \alpha_{ij} \mathbb{1}_{\text{DayType}_t=i} \mathbb{1}_{\text{Offset}_t=j} + s_1(\text{Toy}_t) \\
 & + g_1(\text{CC}_t) + g_2(T_t) + g_3(W_t) \\
 & + k_1(\theta_t) + k_2(\theta_t^{\text{Min}}) + k_3(\theta_t^{\text{Max}}) + \sum_{i=1}^{12} \beta_i T_{t-24} \mathbb{1}_{\text{Month}_t=i} + \varepsilon_t,
 \end{aligned} \tag{2.1}$$

où

- Y_t est la demande en électricité pour l'observation t
- Modélisation de la tendance :
 - t est le nombre d'observations depuis le début de l'échantillon

- $Theating_t$ est le nombre de données depuis le début de la base de données multiplié par la température
- Modélisation des saisonnalités et données calendaires :
 - $DayType_t$ indique le type de jours pour l'observation t
 - $Month_t$ indique le mois de l'observation t
 - $Offset_t$ indique si l'observation t est en heure d'été ou en heure d'hiver
 - Toy_t est le moment dans l'année de l'observation t
- Modélisation de la thermosensibilité en utilisant des variables météorologiques brutes :
 - CC_t est la nébulosité pour l'observation t
 - T_t est la température pour l'observation t
 - T_{t-24} est la température retardée de 24 heures
 - W_t est la vitesse du vent pour l'observation t
- Modélisation de la thermosensibilité en utilisant des températures lissées :
 - θ_t est un lissage exponentiel de la température observée T_t pour l'observation t : $\theta_t = (1 - \alpha)T_t + \alpha\theta_{t-1}$ où $\alpha \in [0, 1]$ contrôle le degré de lissage
 - θ_t^{Min} et θ_t^{Max} sont les températures maximale et minimale journalières lissées
- ε_t bruit centré supposé Gaussien

La covariable *Theating* représente la tendance chauffage. Elle sert à modéliser l'impact des changements économiques sur le chauffage ainsi que l'amélioration de l'efficacité thermique des bâtiments modernes et des technologies de chauffage. La tendance pour la climatisation n'est pas significative.

La variable *DayType* prend pour valeurs 1 le dimanche, 2 le lundi, 3 les mardi-mercredi-jeudi, 4 le vendredi, 5 le samedi, et 6 et 7 les dimanches respectivement de décembre et de juillet. Les vacances ne sont pas considérées dans les modèles. Classiquement, les variables *Offset* et *Toy* servent à modéliser respectivement la rupture due au changement d'heure et la saisonnalité annuelle.

La variabilité thermique importante en inter-saison est prise en compte par la variable T_{t-24} . Les variables météorologiques agrégées ne sont pas sélectionnées.

Ce modèle est obtenu par une procédure longue, fastidieuse et surtout ne pouvant pas être généralisée automatiquement à d'autres jeux de données, même similaires au portefeuille EDF. Ceci motive l'utilisation des procédures automatiques de type Post2.

2.1.4 Modèles additifs moyen terme construits avec une sélection automatisée

Les variations de la dimension du portefeuille et de l'activité socio-économique induisent des phénomènes basse fréquence (tendances) qui ont un impact important sur la qualité de prévision. Nous nous affranchissons de ce problème, qui est en dehors de notre étude, en estimant le modèle de l'expert sur l'ensemble de la période d'étude (échantillons d'apprentissage et de test) puis en retranchant les effets modélisant la tendance dans le modèle (2.1) à la charge consommée (voir Annexe C.3.1). Dans la suite de ce document, nous ne travaillons que sur ces données corrigées de la tendance. Nous avons aussi travaillé sur les données non corrigées, avec des conclusions similaires à celles présentées dans le document. Appliqué à ce nouveau jeu de données, le modèle EDF est le modèle (2.1) auquel les parties modélisant la tendance sont retirées.

2.1.4.1 Sélection automatique de variables

Du fait des propriétés du Group LASSO, les procédures ne réalisent pas de la sélection simultanée de variables continues et factorielles de manière optimale. Comme nous nous intéressons plus particulièrement à la modélisation de la part climatique, nous conservons un modèle qui n'est pas remis en cause et qui est présent dans chaque modèle sélectionné. Ce modèle, appelé modèle

permanent, est donné par l'équation (2.2).

$$Y_t = \beta_0 + \sum_{i=1}^7 \sum_{j=1}^3 \alpha_{ij} \mathbb{1}_{DayType_t=i} \mathbb{1}_{Offset_t=j} + \sum_{i=1}^{12} \beta_i T_{t-12} \mathbb{1}_{Month_t=i} + s_1(Toy_t) + \varepsilon_t. \quad (2.2)$$

Le dictionnaire de covariables candidates, qui contient les variables météorologiques, est de taille 76. Nous appliquons Post2 pour sélectionner des variables dans ce dictionnaire.

Puisque nous travaillons demi-heure par demi-heure, nous sélectionnons un sous-ensemble de covariables différent selon l'instant de la journée.

Similitudes et différences entre les variables sélectionnées pour le modèle EDF et par les procédures automatiques La Figure 2.5 présente des covariables sélectionnées par la procédure Post2Gcv selon l'instant (les 76 covariables ne sont pas représentées par souci de lisibilité). Sur cette figure, pour un instant et une covariable donnés, il y a un point lorsque la covariable fait partie du sous-ensemble de covariables sélectionnées à cet instant. Inversement, l'absence d'objet indique que la covariable n'est pas sélectionnée pour cet instant. Par exemple, la nébulosité a été sélectionnée à l'instant 24, mais pas à l'instant 8.

Les procédures et la sélection de l'expert ont des résultats communs. Par exemple, comme pour le modèle EDF, la variable nébulosité est préférée à la variable nébulosité tronquée pour les températures chaudes. La nébulosité influence le chauffage et l'éclairage. Indépendamment de la température, une nébulosité forte conduit à un besoin plus important d'éclairage le jour. Il y a évidemment des différences. L'expert sélectionne un seul facteur de lissage par variable lissée, alors que nos procédures peuvent en sélectionner plusieurs. Les facteurs de lissage les plus souvent sélectionnés par la procédure Post2Gcv sont différents de ceux utilisés pour le modèle EDF. Cependant, ce sont les mêmes covariables dont des lissages sont sélectionnés (températures brute, ainsi que maximale et minimale journalières).

Cohérence physique de la sélection automatique Nous retrouvons certains résultats intuitifs. Par exemple, la nébulosité est principalement sélectionnée le jour. Celle-ci influence principalement l'utilisation de l'éclairage. Naturellement, pendant la journée, lorsque le ciel est très couvert, les clients utilisent plus de lumière. Par contre, la nébulosité influence peu l'utilisation de l'éclairage la nuit. Les températures maximale et minimale journalières sont principalement sélectionnées en soirée et respectivement autour de midi et le matin. Ceci est cohérent, car cela combine l'inertie des bâtiments avec le fait que la température maximale, par exemple, soit généralement atteinte en début d'après-midi.

Comparaison entre les différentes procédures automatiques Nous comparons ensuite les procédures de sélection Post2Bic, Post2Aic et Post2Gcv. La Figure 2.6, qui se lit comme la Figure 2.5, récapitule des variables sélectionnées par les trois procédures.

Comme les critères AIC et GCV, fondés sur des minimisations de distances, ont des propriétés proches, les sous-ensembles de covariables sélectionnées par Post2Aic et Post2Gcv sont très proches. Post2Bic sélectionne moins de covariables, illustrant ainsi la plus grande sévérité du BIC. Les modèles ainsi sélectionnés sont plus proches du modèle EDF, car il tend à identifier peu de covariables par type de variables (par exemple, une ou deux températures lissées, etc). Il sélectionne des lissages proches de ceux du modèle EDF pour les températures maximale et minimale, ainsi que principalement un coefficient faible et un proche de celui du modèle EDF pour le lissage de la température.

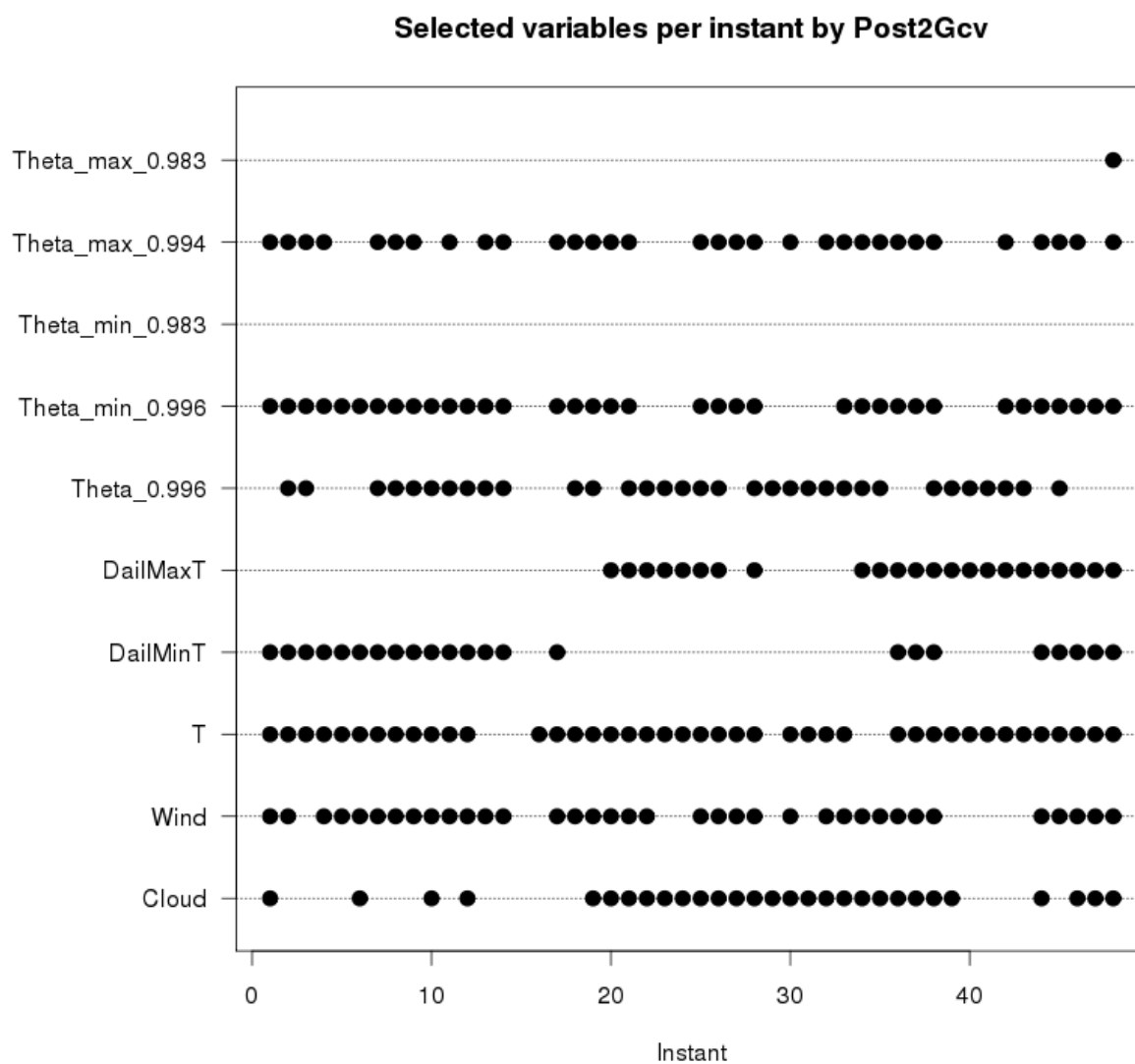


FIGURE 2.5 – Quelques covariables sélectionnées par Post2Gcv selon l'instant

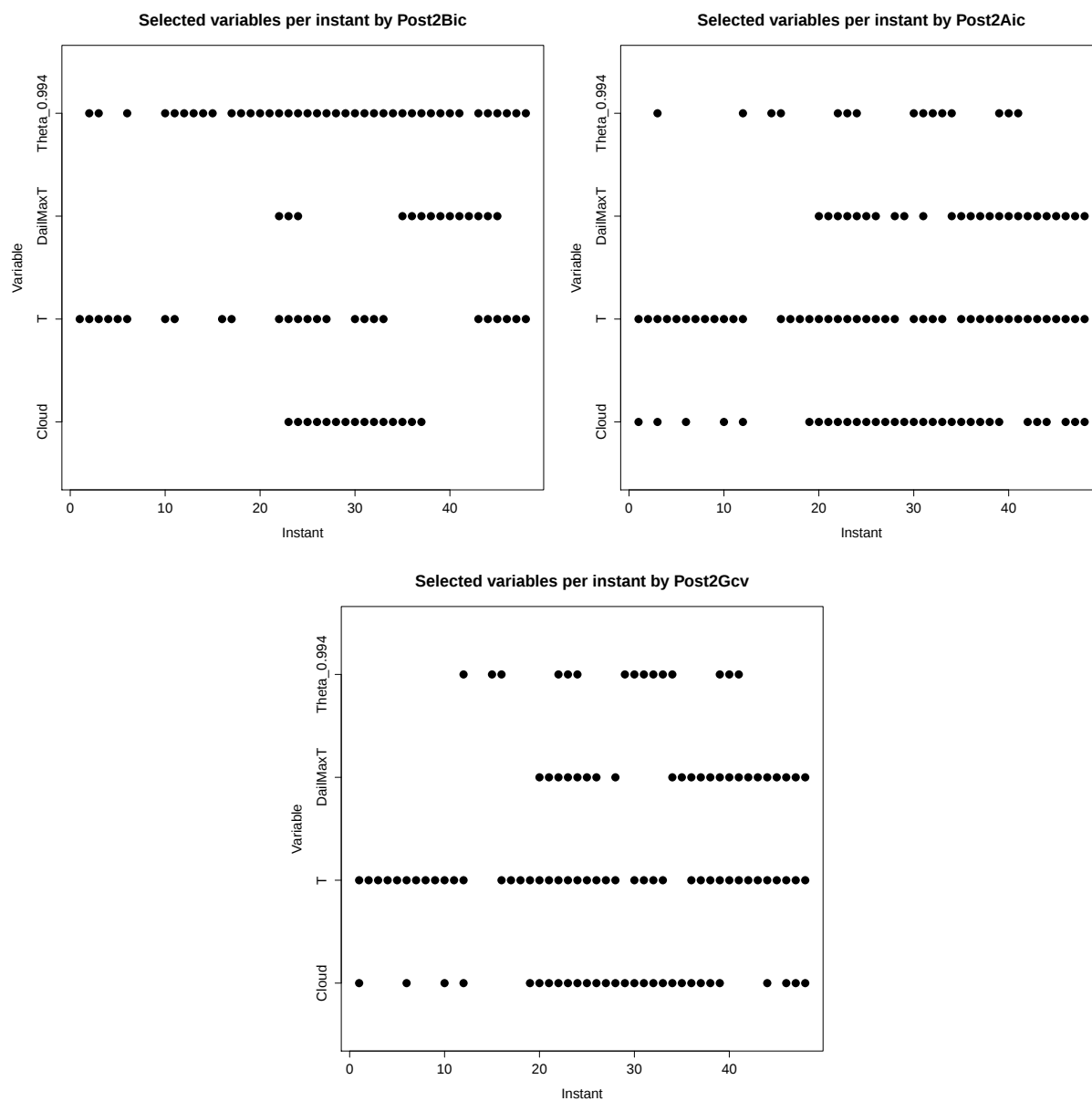


FIGURE 2.6 – Quelques covariables sélectionnées par Post2Bic, Post2Aic et Post2Gcv selon l'instant de la journée

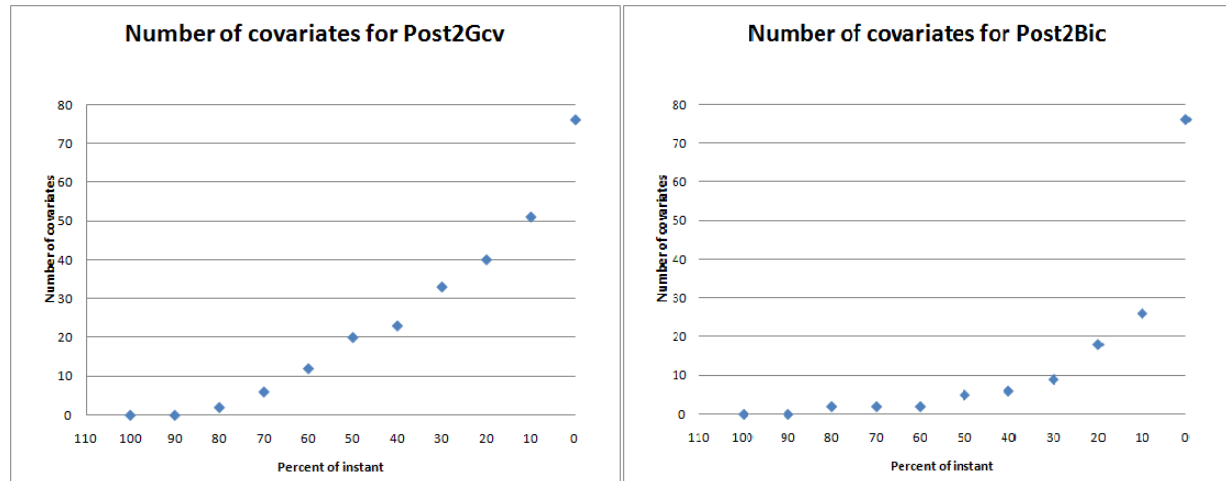


FIGURE 2.7 – Nombre de covariables conservées pour au moins un pourcentage d’instant pour Post2Gcv et Post2Bic

Post-traitement des sous-ensembles sélectionnés En construisant un modèle différent selon l’instant, nous supposons l’indépendance entre chaque instant. Ajouter a posteriori l’information que les instants sont dépendants peut améliorer les performances. Nous présentons un type de post-traitement dans ce chapitre. Un second est présenté en Annexe C.3.2.

Le post-traitement, caractérisé de “global”, consiste à ne conserver qu’un unique sous-ensemble de covariables commun à tous les instants. La Figure 2.7 présente le nombre de covariables sélectionnées pour au moins un certain pourcentage des instants. Par exemple, deux covariables ont été conservées pour au moins 80% des instants pour Post2Gcv.

Nous initialisons une grille de pourcentages de présence des covariables (voir les abscisses de la Figure 2.8) et appliquons ensuite une technique du “coude” [110] sur le RMSE (voir Annexe C.2) en estimation des modèles obtenus pour sélectionner les modèles post-traités. En plus d’être simple à implémenter et à automatiser, le post-traitement “global” permet d’homogénéiser les sous-ensembles de covariables utilisées pour modéliser la charge consommée. Dans le package **CASA**, nous proposons aussi un post-traitement entièrement automatisé permettant à l’utilisateur de sélectionner le sous-ensemble de covariables minimisant un critère de sélection voulu sur un échantillon de validation. Les défauts du post-traitement sont la perte de flexibilité à laquelle il conduit et le risque de sélectionner un nombre inadéquat de variables.

La technique du “coude” conduit à sélectionner les modèles comportant les covariables présentes dans au moins 30% et 20% des instants pour Post2Gcv et pour Post2Bic respectivement. Nous notons ces modèles $\text{Post2Gcv} > 0.3$ et $\text{Post2Bic} > 0.2$. Sélectionner $\text{Post2Gcv} > 0.2$ et $\text{Post2Bic} > 0.1$ n’améliore pas fortement le RMSE, mais complique fortement les modèles, alors que $\text{Post2Gcv} > 0.3$ et $\text{Post2Bic} > 0.2$ permettent de fortement diminuer le RMSE par rapport à respectivement $\text{Post2Gcv} > 0.4$ et $\text{Post2Bic} > 0.3$. Ces deux modèles comportent respectivement 33 et 18 covariables. Nous présentons les détails du modèle $\text{Post2Bic} > 0.2$ en Annexe C.3.3.

La Table 2.1 récapitule le nom des modèles sélectionnés, ainsi que quelques détails. Ensuite, nous testons ces modèles en prévision sur la période septembre 2012-juillet 2013.

2.1.4.2 Performances prédictives

En plus du modèle EDF optimisé sur les données agrégées au niveau France, nous introduisons un benchmark supplémentaire : le modèle additif proposé par Goude, Nedellec et Kong [55] comparant ainsi le modèle EDF et les modèles issus de nos procédures avec un modèle additif

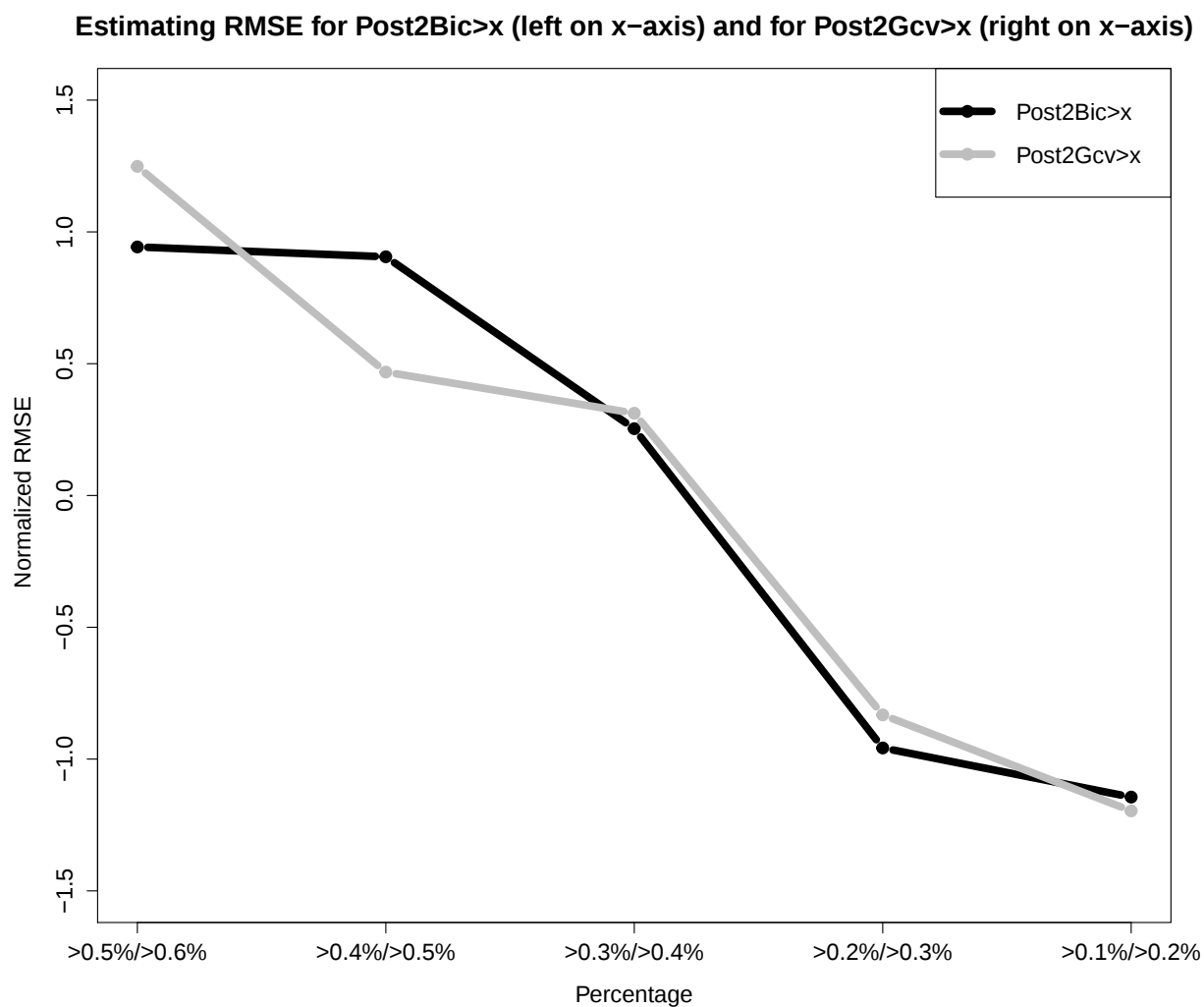


FIGURE 2.8 – RMSE en estimation selon le pourcentage considéré

Modèle	Post-traitement des sous-ensembles de covariables	Nombre de sous-ensembles	Détails
EDF (ou Benchmark)	Non	1 commun à tous les instants	Recherche manuelle
Post2Bic	Non	48 (1 par instant)	Recherche automatique
Post2Aic	Non	48 (1 par instant)	Recherche automatique
Post2Gcv	Non	48 (1 par instant)	Recherche automatique
Post2Bic>0.2	Oui	1 commun à tous les instants	Post-traitement “global”
Post2Gcv>0.3	Oui	1 commun à tous les instants	Post-traitement “global”

TABLE 2.1 – Modèles sélectionnés

générique.

Un modèle de prévision de consommation classique : Goude *et al.* [55] Le modèle BenchMT1 s’inspire du modèle proposé par Goude *et al.* [55] qui a été optimisé pour les postes sources.

$$\begin{aligned}
 Y_t = & \beta_0 + f_1(T_t) + f_2(T_{t-48}) + f_3(T_{t-96}) + f_4(\theta_t) \\
 & + f_5(Toy_t) + \sum_{i=1}^7 \alpha_i \mathbb{1}_{DayType_t=i} + \beta_1 \mathbb{1}_{Offset_t=1} + \gamma \mathbb{1}_{SpecialTarif_t=1} \\
 & + \varepsilon_t,
 \end{aligned}$$

où

- θ_t température lissée de facteur de lissage 0.99
- $SpecialTarif_t = 1$ si il y un EJP en cours, 0 sinon

Ce modèle moyen terme générique a obtenu des bonnes performances sur les données locales.

Comparaison des performances des différentes procédures La Table 2.2 fournit les critères d’évaluation de la performance en prévision des procédures (voir Annexe C.2). La Figure 2.9 donne le rapport entre le RMSE en prévision de certaines procédures et celui du modèle EDF. Si ce rapport est supérieur à 1, alors le RMSE de la procédure en question est supérieur à celui du modèle EDF, et donc celui-ci a, selon ce critère, de moins bonnes performances prédictives.

	MAPE	MAE	RMSE
Post2Bic>0.2	1.12%	505	645
Post2Gcv>0.3	1.15%	512	648
Post2Aic	1.17%	523	663
Modèle EDF	1.16%	519	667
Post2Gcv	1.17%	526	667
Post2Bic	1.24%	562	730

TABLE 2.2 – Critère de performance en prévision

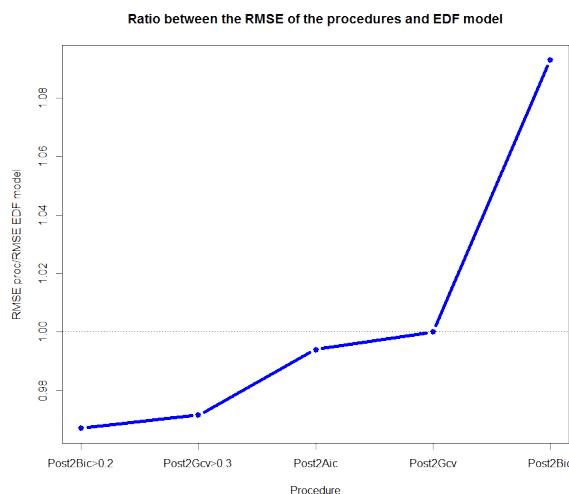


FIGURE 2.9 – RMSE en prévision des procédures nommées en abscisse divisé par celui du modèle EDF

Les performances de Post2Aic et Post2Gcv, dont la phase de sélection a été automatique, sont équivalentes à celles du modèle EDF, construit manuellement. Après le post-traitement des ensembles de covariables sélectionnées, les performances sont meilleures. Notons que le modèle BenchMT1 conduit à un MAPE de 2.00% et à un RMSE de 1173. Les modèles obtenus à l’aide de la sélection de variables (automatique ou manuelle) ont donc des meilleures performances qu’un modèle additif générique.

Post2Aic a des meilleures performances que Post2Bic, illustrant ainsi les bonnes propriétés prédictives de l’AIC. Le post-traitement “global” améliore les performances, notamment pour Post2Bic>0.2. Nous le justifions par le fait que le BIC identifie les covariables les plus influentes, mais à cause de sa sévérité, ne les sélectionne pas pour suffisamment d’instant. Les variables

présentes dans $\text{Post2Bic} > 0.2$ ont été sélectionnées pour au moins 10 instants (c'est-à-dire pour 5 heures). Ce pas relativement fin permet de capter des effets limités dans la journée (typiquement la nébulosité). En termes de RMSE, le modèle EDF a des moins bonnes performances prédictives, excepté par rapport à Post2Bic et à Post2Gcv . En termes de MAPE et de MAE, seuls $\text{Post2Bic} > 0.2$ et $\text{Post2Gcv} > 0.3$ sont plus performants. Le modèle EDF a donc vraisemblablement des erreurs extrêmes plus fortes que les autres procédures (sauf Post2Bic). Nous confirmons ce propos sur la Figure C.6 de l'Annexe C.3.4, qui présente le rapport entre les valeurs absolues par ordre croissant des erreurs de $\text{Post2Bic} > 0.2$ et de Post2Gcv avec celles du modèle EDF. Les erreurs plus élevées du modèle EDF peuvent provenir des événements météorologiques extrêmes sur lesquels il a moins de paramètres pour agir.

Pour le post-traitement "global", nous avons sélectionné $\text{Post2Bic} > 0.2$ et $\text{Post2Gcv} > 0.3$. La Figure 2.10 justifie a posteriori ce choix. Le RMSE de chaque procédure est représenté sur cette figure.

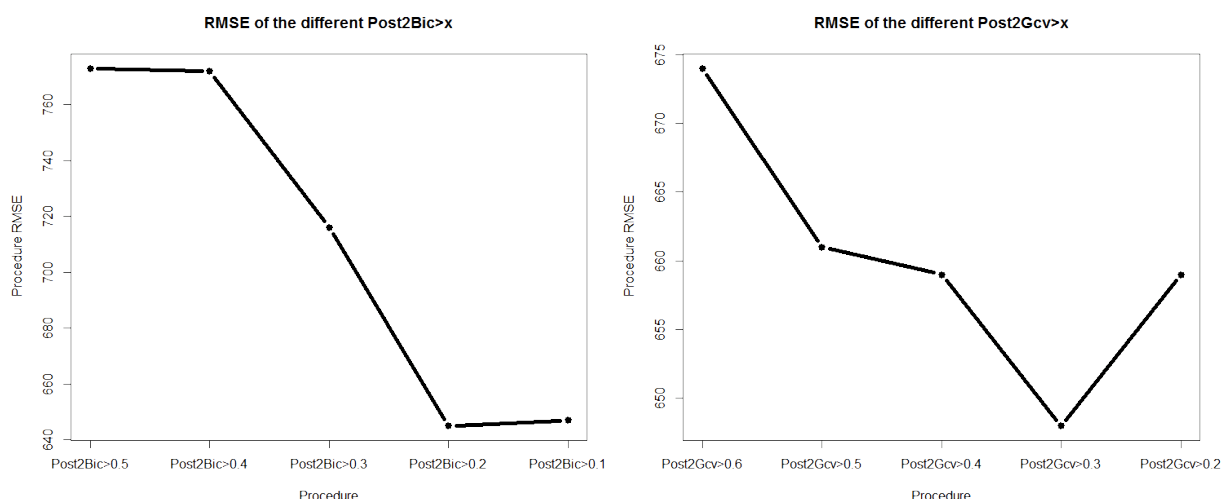


FIGURE 2.10 – RMSE en prévision selon le pourcentage considéré

Les modèles $\text{Post2Gcv} > 0.6$, $\text{Post2Gcv} > 0.5$, $\text{Post2Gcv} > 0.4$ ainsi que $\text{Post2Bic} > 0.5$, $\text{Post2Bic} > 0.4$, $\text{Post2Bic} > 0.3$ sont sous-déterminés : il manque des covariables influentes. Inversement, $\text{Post2Gcv} > 0.2$ et $\text{Post2Bic} > 0.1$ sur-apprennent l'échantillon d'apprentissage.

Ensuite, nous justifions l'intérêt d'introduire un correctif court terme en étudiant les performances et les erreurs de Post2Gcv . Nous développons ce modèle, bien qu'il n'ait pas les meilleures performances, car il est issu d'une procédure Post2 sans post-traitement. Ses performances sont très proches de celles du modèle EDF, ce qui permet d'établir un parallèle entre les deux. Enfin, le GCV est le critère de sélection utilisé par l'expert.

L'étude des erreurs est détaillée en Annexe C.3.4. Nous retrouvons les mêmes types d'erreurs pour Post2Gcv et le modèle EDF. Les performances sont moins bonnes en novembre 2012 et en avril 2013. La température, qui a été anormalement froide durant le second mois, notamment autour de 8-9h, influence fortement les erreurs. Pour les autres mois de l'échantillon test, les variables météorologiques ont peu d'influence sur les erreurs des modèles. Celles-ci ne sont pas centrées autour de 0, ce que nous expliquons par une tendance corrigée de manière incomplète. Nous nous intéressons ensuite au PACF (fonction d'auto-corrélation partielle) des erreurs en estimation et en prévision pour justifier l'intérêt d'introduire une correction court terme des modèles.

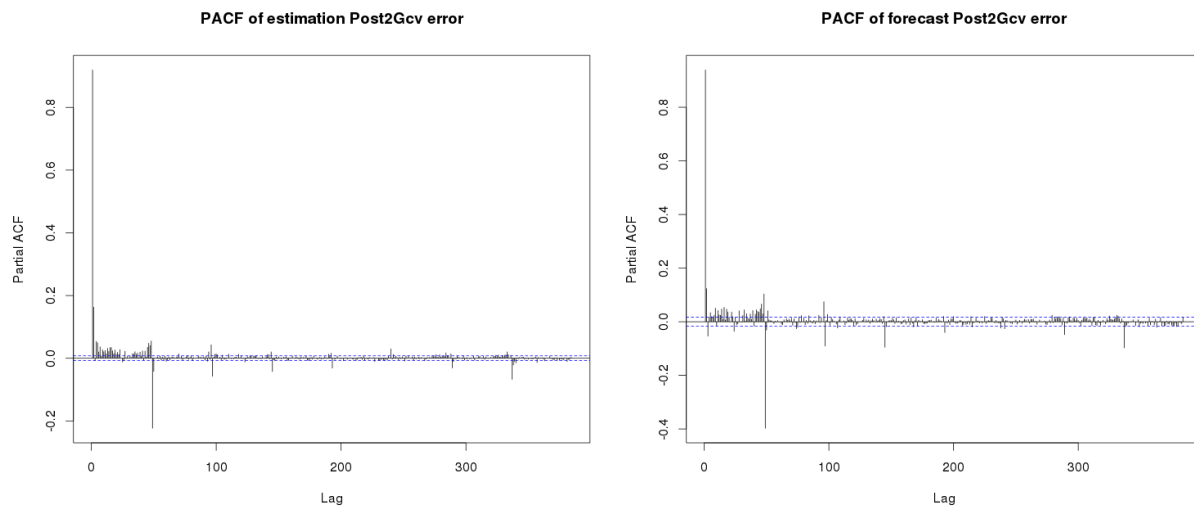


FIGURE 2.11 – PACF en estimation (gauche) et en prévision (droite) pour Post2Gcv

L'aspect séries temporelles de la charge consommée n'a pas été exploité à moyen terme, que ce soit sur les échantillons d'apprentissage ou de test. Pour un tel horizon, les charges consommées dans un passé récent sont inconnues. Classiquement, l'erreur commise pour une observation est plus influencée par celle commise la veille et lorsque le type de jours correspond que par celle commise respectivement 12 heures auparavant (par exemple) et lorsque le type de jours ne correspond pas, même lorsque celle-ci correspond à une observation plus proche.

2.1.5 Modèles court terme

Les erreurs d'un modèle moyen terme ont plusieurs origines. La modélisation peut être imparfaite, avec par exemple l'absence de variables influentes, un choix de famille de modèles qui approximent mal la réalité ou des effets estimés dans des bases non optimales. La non-stationnarité des séries due à un événement ponctuel conduit à des erreurs caractérisées de haute fréquence. Nous nous servons des erreurs passées pour corriger ces deux types d'erreurs.

Dans le paragraphe suivant, nous présentons deux méthodes possibles pour estimer un modèle court terme à un horizon de 24 heures.

2.1.5.1 Présentation des méthodes court terme

Différentes approches ont été proposées dans la littérature pour la prévision court terme : par exemple, en ajoutant les charges retardées dans l'équation moyen terme (voir [94], [46], [17]), en utilisant des procédures en plusieurs étapes comme décrites dans [91], ou encore en considérant des bruits auto-regressifs dans l'équation du moyen terme comme [120]. Nous présentons ensuite les deux premières.

Ajout de charges retardées dans le modèle moyen terme Le modèle de Pierrot et Goude [94], que nous notons BenchST2, suit cette philosophie. Goude *et al.* [55] présentent un autre modèle de ce type, qui inspire le modèle suivant, nommé BenchST1.

$$\begin{aligned}
Y_t = & \beta_0 + f_1(T_t) + f_2(T_{t-48}) + f_3(T_{t-96}) + f_4(\theta_t) + f_5(T_{oy_t}) \\
& + \sum_{i=1}^7 \alpha_i \mathbb{1}_{DayType_i=i} + \beta_1 \mathbb{1}_{Offset=1} + \gamma \mathbb{1}_{SpecialTarif=1} \\
& + \mathbf{m}(\mathbf{Y}_{t-48}) + \varepsilon_t,
\end{aligned}$$

Comme Pierrot et Goude [94], Goude *et al.* [55] modélisent l'effet de la charge retardée par le biais d'une fonction lisse. Ce modèle se décompose en une partie moyen terme, comportant des températures brutes (retardées ou pas), une température lissée, la position dans l'année, le type de jours, des offsets et des données de tarification, et en une partie court terme, composée de la charge retardée de 24 heures. Avec cette méthodologie de modélisation, il est plus difficile d'interpréter les fonctions estimées pour la partie moyen terme, car l'introduction de la charge retardée de 24 heures apporte beaucoup d'informations, notamment propres à la thermosensibilité par exemple.

La seconde méthode consiste à considérer que nous captons les effets basse fréquence grâce à un modèle moyen terme. Nous ajoutons alors l'information des erreurs passées pour corriger ce modèle et ainsi capter des effets haute fréquence.

Correctif court terme Pour cette correction, nous supposons qu'un modèle moyen terme (noté MT) a été estimé sur une période d'apprentissage. Nous l'utilisons pendant quelques mois, voire un an, à un horizon h , c'est-à-dire que si nous voulons prévoir la charge consommée en t , nous avons la connaissance exhaustive des informations jusqu'à l'observation $t-h$. Nous utilisons ces informations pour corriger la prévision du modèle MT en fonction de ses erreurs passées (voir Figure 2.12). Nous notons $\hat{f}^{MT}$ la fonction de régression estimée pour le modèle MT. Nous résumons ensuite le principe de correctif court terme.

1. Supposons que la charge consommée Y_t est modélisée par le modèle suivant : $Y_t = f(\mathbf{X}_t) + \varepsilon_t$, où $\mathbf{X}_t$ covariables. Notons $\hat{f}$ estimateur de f estimé sur un échantillon d'apprentissage, alors $\hat{Y}_t = \hat{f}(\mathbf{X}_t)$ et $\hat{\varepsilon}_t = Y_t - \hat{Y}_t$
2. Supposons que $\hat{\varepsilon}_t$ est modélisée par le processus auto-régressif suivant : $\hat{\varepsilon}_t = g((\hat{\varepsilon}_{t-i})_{i \geq h}) + \nu_t$, où ν_t bruit blanc. Notons $\hat{g}$ estimateur de g estimé dans le passé (récent), alors $\bar{\varepsilon}_t = \hat{g}((\hat{\varepsilon}_{t-i})_{i \geq h})$
3. $\hat{Y}_t^{CT} = \hat{Y}_t + \bar{\varepsilon}_t$

La Figure 2.12 explique le principe de correction court terme. Sur cette figure, les charges consommées (rouge) et prédites par un modèle moyen terme (trait bleu) sont connues avant l'observation t . Nous avons donc accès aux erreurs passées du modèle. À partir de l'instant t , nous connaissons uniquement les prévisions futures du modèle (pointillé bleu). Nous utilisons les erreurs passées du modèle pour prévoir ses erreurs futures puis nous corrigeons ses prévisions futures avec ces informations pour prévoir la consommation future (rose).

Nous utilisons la seconde méthode pour des raisons métiers (elle est plus utilisée à EDF ; l'information fournie au modèle étant différente, elle s'est avérée plus performante pour un horizon d'étude de 24 heures sur d'autres jeux de données et l'autre méthode rend compliquée l'interprétation des effets) et de cohérence par rapport à notre étude, puisque nous poursuivons l'étude des modèles moyen terme présentés dans la section précédente et vérifions que le court terme préserve les mêmes propriétés de performance. Nous comparons nos modèles moyen terme corrigés avec BenchST1 et BenchST2. Nos correctifs utilisent les erreurs de modèles moyen terme, alors que BenchST1 et BenchST2 utilisent les charges retardées. Il est plus informatif d'utiliser les erreurs d'un modèle moyen terme lorsque celui est bon, puisque cela donne une information supplémentaire sur la complexité des prévisions récentes.

Principe de la prévision court terme

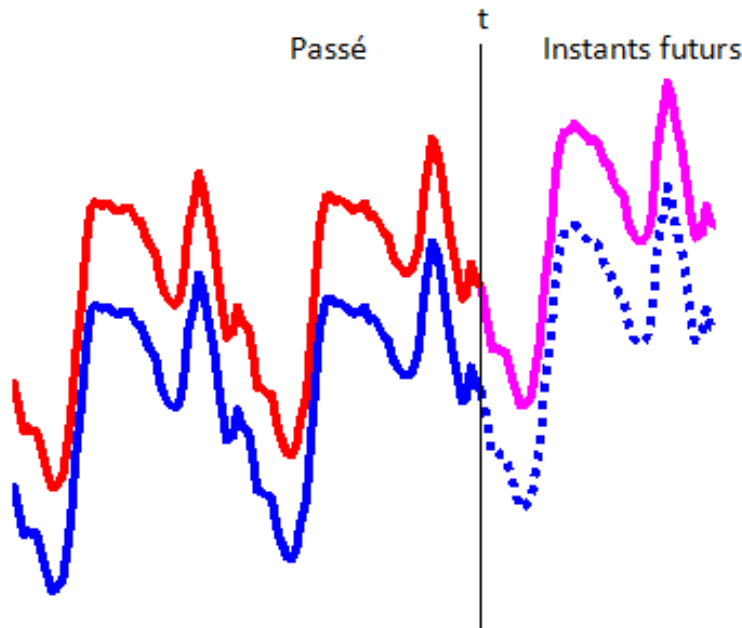


FIGURE 2.12 – Principe de la correction court terme : le trait bleu représente les prévisions passées du modèle (connues), les pointillés bleus les prévisions futures (connues), le trait rouge les vraies charges consommées passées (connues) et le trait rose les vraies charges consommées futures (inconnues)

À l'horizon d'un jour, nous constatons sur le PACF que les erreurs moyen terme ont une forte corrélation partielle avec celles de la veille, celles de la veille plus une demi-heure, celles de l'avant veille, celles de trois jours avant et celles d'une semaine avant. D'autres sources d'erreurs sont la présence de biais éventuels introduits par le modèle moyen terme et une mauvaise extrapolation d'effets basse fréquence tels que la tendance. Ces observations nous conduisent à considérer les modèles (2.3) et (2.4).

— Correction linéaire :

$$\hat{\varepsilon}_t = \sum_{i=1}^3 \beta_i \hat{\varepsilon}_{t-48 \times i} + \beta_4 \hat{\varepsilon}_{t-49} + \beta_5 \hat{\varepsilon}_{t-336} + u_t. \quad (2.3)$$

— Double correction :

$$\hat{\varepsilon}_t = \sum_{i=1}^3 \beta_i \hat{\varepsilon}_{t-48 \times i} + \beta_4 \hat{\varepsilon}_{t-49} + \beta_5 \hat{\varepsilon}_{t-336} + \mu_t + u_t, \quad (2.4)$$

où u_t est un bruit blanc et μ_t un processus de moyenne mobile. Nous commençons par estimer les $(\beta_i)_i$ sur l'échantillon d'apprentissage du modèle moyen terme, μ_t est alors considéré comme une constante, qui est mise à jour en ligne durant le processus de prévision court terme. L'intérêt de l'introduction de la moyenne mobile est de lisser les éventuels écarts accidentels. Elle corrige des biais qui évolueraient lentement lors du processus de prévision.

2.1.5.2 Performances court terme

Ici, l'introduction d'une moyenne mobile n'améliore pas les prévisions. Par contre, elle les améliore sur les données où la tendance n'a pas été corrigée au préalable. Le passage au court

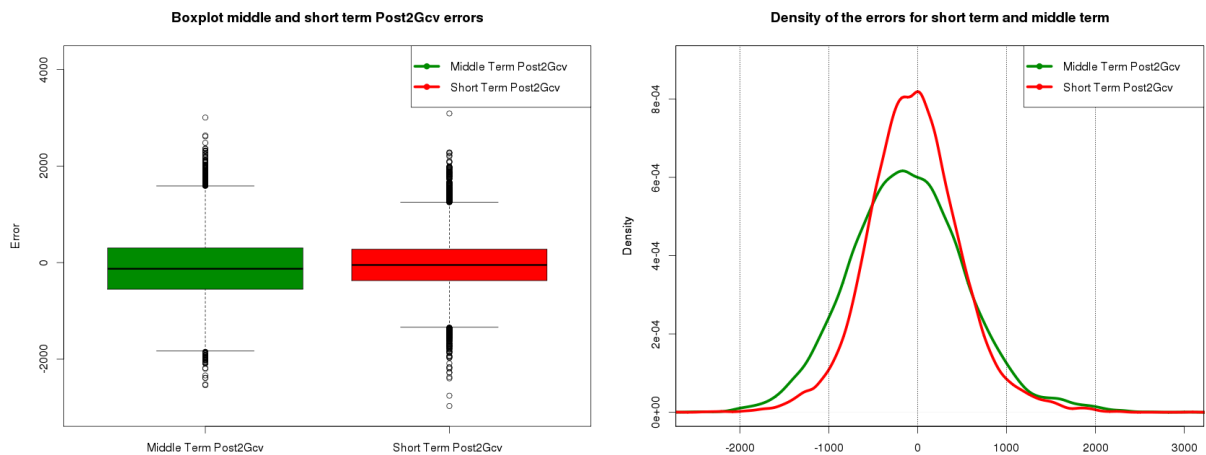


FIGURE 2.13 – Erreurs en prévision moyen et court termes de Post2Gcv

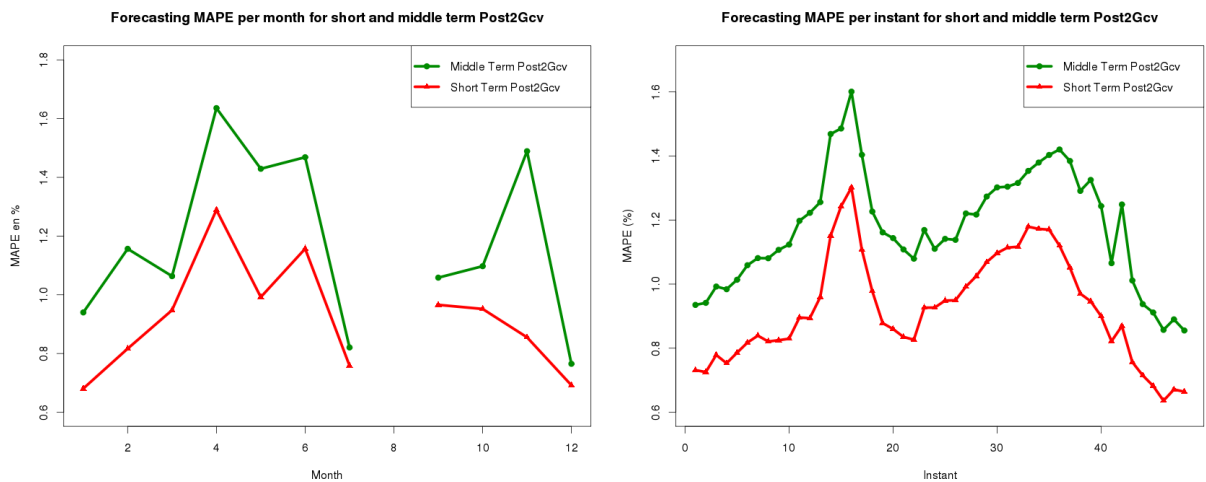


FIGURE 2.14 – Erreurs en prévision par mois (gauche) et selon l'instant (droite)

terme améliore les performances en prévision des modèles moyen terme. Ceci est la conséquence de la correction du biais et de la réduction de la variance des erreurs, comme le montre la Figure 2.13 qui fournit les boîtes à moustaches et une estimation des densités des erreurs court et moyen termes de Post2Gcv .

La Figure 2.14 représente le MAPE en prévision par mois (gauche) et selon l'instant (droite) pour Post2Gcv. La correction court terme améliore les performances pour tous les mois. Le mois de novembre, contrairement au mois d'avril, est bien corrigé, ce qui est cohérent avec les origines supposées des erreurs. Tous les instants sont corrigés, et notamment le pic d'erreurs du soir, qui est un point critique pour les producteurs et les fournisseurs d'électricité en France.

La Table 2.3 fournit le MAPE et le RMSE des procédures corrigées avec (2.3). Nous testons aussi une méthode de mélange de prédicteurs, qui est un sujet important à EDF, car nous avons estimé beaucoup de modèles pouvant servir d'experts. Aggregation est obtenue en agrégeant tous les experts issus de Post2 corrigés par (2.3) et (2.4) grâce à la fonction *ewa* du package **OPERA** (voir [50]). La prévision naïve consiste à prévoir la charge consommée par la charge consommée la veille.

Version corrigée du modèle moyen terme	MAPE	RMSE
Aggregation	0.89	518
Post2Gcv>0.3	0.91	526
Post2Aic	0.92	529
Post2Bic>0.2	0.92	531
Post2Gcv	0.92	532
EDF model	0.93	537
Post2Bic	0.98	571
Bench ST2	1.49	899
Bench ST1	1.51	901
Naïve	5.88	3889

TABLE 2.3 – MAPE et RMSE en prévision pour le court terme

Toutes les procédures étudiées obtiennent de très bonnes performances avec des MAPE inférieurs à 1%. Les écarts de performances se réduisent sur des prévisions court terme. Cependant, même si l'amélioration est faible, les procédures issues de Post2 ont de meilleures performances que le modèle EDF corrigé. La comparaison des performances reste similaire au court terme par rapport au moyen terme. Le post-traitement des sous-ensembles de covariables sélectionnées améliore les performances. Malgré un faible gain, la procédure Post2 est une procédure entièrement automatique demandant beaucoup moins d'interventions humaines et de connaissances métiers que la procédure utilisée pour construire le modèle EDF. L'agrégation des experts issus de la procédure Post2 offre une solution améliorant les performances en prévision.

2.2 Modélisation et prévision de consommation pour des mailles locales de réseaux de distribution

La prise en compte de la météorologie est souvent cruciale pour améliorer la prévision de consommation électrique. Un des aspects de la modélisation climatique concerne la localisation spatiale de sa mesure. Cette information est liée à la topologie du réseau qui est complexe, évolutive et parfois indisponible. Revenons un instant au portefeuille EDF. Pour celui-ci, des moyennes pondérées nationales sont utilisées [21], le problème statistique de sélection consistant alors en la sélection de versions modifiées de la température, de la nébulosité ou de la vitesse du vent. Cependant, avec le développement des réseaux intelligents, la sélection de la localisation de la mesure météorologique devient une problématique à fort enjeu sur laquelle peu de travaux ont été réalisés. Ici, l'utilisation de moyennes nationales n'est pas optimale. Le problème de sélection associé à ces réseaux locaux est plus compliqué puisqu'en plus de la sélection des covariables météorologiques, la localisation de leurs mesures est aussi problématique. Nous nous concentrons dans cette section à ce problème de localisation des mesures. Comme il y a beaucoup de noeuds à prévoir sur un réseau (par exemple, environ 2200 postes sources), le statisticien ne peut pas sélectionner manuellement la ou les stations météorologiques permettant les meilleures prévisions. La sélection des endroits où sont mesurées les variables météorologiques pose les deux questions suivantes :

- Combien de stations météorologiques utiliser ?
- Quelles stations météorologiques ?

Beaucoup de travaux traitant de la sélection des stations météorologiques dans le cadre de la prévision d'une maille locale de consommation électrique proposent l'utilisation d'un nombre fixé de stations météorologiques pour prévoir la charge consommée pour un noeud de la maille étudiée. Par exemple, Goude *et al.* [55] proposent l'utilisation d'une seule station météorologique pour prévoir la charge consommée au niveau des postes sources (environ 2200 en France). Dans le cadre de la compétition GEFCom 2012, les quatre équipes les mieux classées posaient aussi comme

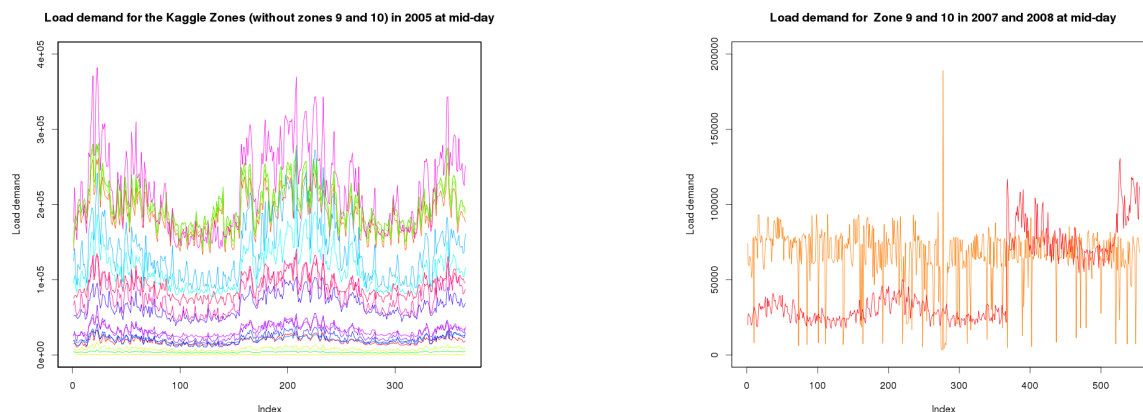


FIGURE 2.15 – Courbes de charge de la compétition GEFCom 2012

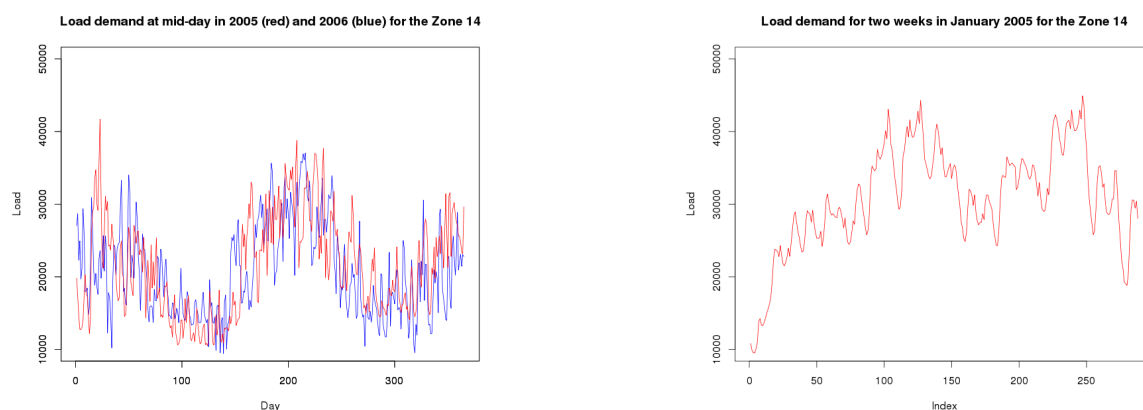


FIGURE 2.16 – Charge consommée en 2005 (rouge) et 2006 (bleue) à midi et deux semaines de charge consommée en janvier 2005 pour la zone 14 de la compétition GEFCom 2012

contrainte un nombre fixe de stations météorologiques (cinq pour Charlton et Singleton [105], onze pour Lloyd [84] et une pour Nedellec *et al.* [91] ainsi que pour Ben Taieb et Hyndman [17]). Cependant, comme Hong *et al.* [63] le remarquent, nous pouvons améliorer les prévisions en relâchant cette contrainte. L'application de nos procédures de sélection automatique répond aux deux questions associées à la sélection de stations météorologiques à la maille locale.

2.2.1 Description des données

Les données de la compétition GEFCom 2012 Les données disponibles des vingt zones à prévoir vont de 2004 à juin 2008 et sont à un pas horaire. Nous apprenons les modèles sur la période 2004-2007 et les testons en prévision en 2008. La Figure 2.15 représente les courbes de charge consommée pour les vingt zones à midi. Les courbes de charge ont une variabilité importante. La consommation est plus forte en hiver et en été. Le niveau des courbes de charge est différent d'une zone à l'autre. Nous excluons de notre étude la zone 7 (identique à la zone 3), la zone 9 (insensible à la température) et la zone 10 (connaissant une forte rupture entre les périodes d'apprentissage et de test). Nous laissons les résultats pour la zone 4, même si la consommation électrique de cette zone est très faible.

La Figure 2.16 fournit la charge consommée à midi en 2005 (rouge) et 2006 (bleue) et pendant deux semaines en janvier 2005 pour la compétition GEFCom 2012. Il y a un cycle journalier. La charge consommée dépend aussi de la position dans l'année. Elle est forte en hiver et en été

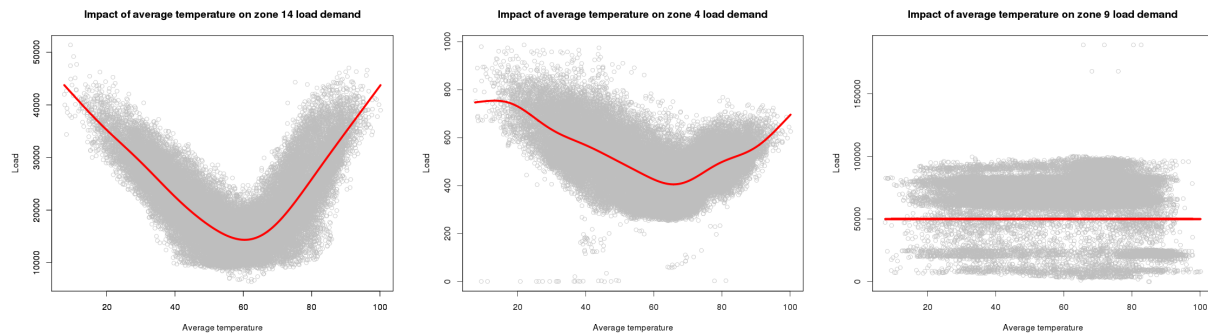


FIGURE 2.17 – Effet de la température moyenne des onze stations météorologiques sur la charge consommée pour trois zones de la compétition GEFCOM 2012

et faible au printemps et en automne. La forte consommation en hiver et en été montre que la température joue un rôle prédominant. La Figure 2.17 donne la charge consommée en fonction de la température moyenne des onze stations météorologiques pour trois zones. La figure montre que pour la première zone, les effets chauffage et climatisation sont élevés. L’effet de la température pour la seconde zone est plus proche de celui observé pour le portefeuille EDF, avec une part chauffage plus forte que la part climatisation. Enfin, la charge consommée pour la dernière zone est insensible à la température.

La Figure 2.18 présente la charge consommée pour la zone 14 en fonction des températures respectivement des stations 4 et 10.

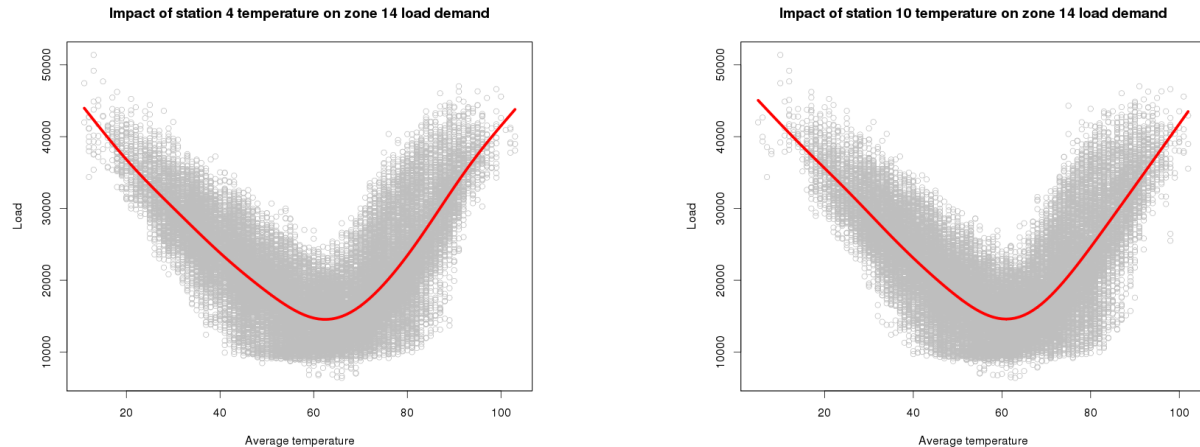


FIGURE 2.18 – Effet de la température des stations météorologiques 4 (gauche) et 10 (droite) sur la charge consommée dans la zone 14 de la compétition Kaggle

Les effets des températures des stations 4 et 10 ont des formes similaires.

Comme pour le portefeuille EDF, la charge consommée ainsi que les effets influençant celle-ci sont fortement dépendants de l’heure de la journée. Nous utilisons donc des modèles différents selon l’heure de la journée.

Les données des postes sources Nous utilisons ces données à un pas horaire. Les données vont du 01/01/2010 au 30/06/2014. Nous apprenons les modèles pendant la période allant de 01/01/2010 au 30/06/2013 et les testons sur la période allant du 01/07/2013 au 30/06/2014. Nous avons accès à des données calendaires, climatiques et tarifaires. Certains postes sources ont

des courbes de charge similaires à celle du portefeuille EDF. Cependant, les profils de clients connectés à un poste source peuvent changer. Par exemple, la charge consommée est en créneaux selon les heures d'activité des industries si le poste source est connecté uniquement à des gros clients industriels. Certains postes sources sont insensibles à la température ou ont des fortes ruptures dans la charge consommée par exemple. Nous utilisons des modèles par instant, qui estiment mieux que les modèles globaux les charges consommées avec un coût informatique plus faible (voir Goude *et al.* [55]).

Différences et similitudes Les deux applications sont locales. Les rapports de signal sur bruit des consommations sont donc plus faibles que celui avec le portefeuille EDF. Les applications sont composées de plusieurs séries temporelles de courbes de charge. Nous disposons des séries de températures de plusieurs stations météorologiques, qui sont mises en compétition pour expliquer les charges consommées. Pour GEFCom 2012, nous n'avons pas d'information géographique aidant au choix de la station, contrairement aux postes sources. Pour ceux-ci, nous pourrions sélectionner la station météorologique la plus proche du poste source au sens d'une distance à déterminer. Cependant, cela pose les trois questions suivantes :

- Pourquoi seulement la plus proche, et non pas les deux ou trois plus proches ?
- Pourquoi la plus proche, alors que certains postes sources, notamment en montagne, ne sont pas nécessairement les plus influencés par la température de la station météorologique la plus proche ?
- Comment définir la notion de plus proche ?

La notion de plus proche nécessite la connaissance topographique du réseau (coordonnées GPS des consommateurs finaux, leur type, leur nombre,...) qui n'est pas toujours disponible. Actuellement, la station météorologique utilisée pour chaque poste source n'a pas été sélectionnée "électriquement". Pour les deux applications, le choix des stations météorologiques demande donc une attention particulière.

Les deux études ont cependant quelques différences. Nous avons potentiellement plus d'informations (localisation, niveau d'agrégation, structure du réseau) pour les postes sources. Contrairement à ces derniers, nous n'avons pas de problème de confidentialité avec les données de GEFCom 2012. Il y a une moins grande diversité des courbes de charge dans la compétition GEFCom 2012 que pour les postes sources.

Nous nous inspirons de modélisations très proches pour les deux applications.

2.2.2 Les modèles de référence

Modèle pour GEFCom 2012 Le modèle de référence que nous utilisons pour traiter des données GEFCom 2012 est inspiré de Nedellec *et al.* [91], qui proposent une modélisation en trois étapes :

1. Modèle long terme : les données sont agrégées sur un mois, un modèle additif modélisant la charge mensuelle par une indicatrice selon le mois et par la température mensuelle de la station météorologique associée à la zone étudiée est estimé. Les résidus de ce modèle sont ensuite lissés avec un noyau Gaussien. Grâce à une interpolation linéaire, une tendance est obtenue. Les charges étudiées sont ensuite corrigées de cette tendance.
2. Modèle moyen terme : présenté par la suite.
3. Modèle court terme : correction du modèle moyen terme, avec par exemple des forêts aléatoires.

Les stations météorologiques sont sélectionnées en minimisant le GCV.

Le modèle de Nedellec *et al.* [91], pour lequel nous ne faisons pas apparaître la dépendance en l'instant de la journée par souci de lisibilité, est le suivant :

$$Y_{t,j}^{det} = \sum_{q=1}^7 m_q \mathbb{1}_{DayType_t=q} + g_1(\theta_{t,k_j}) + g_2(T_{t,k_j}) \\ + g_3(T_{t-1,k_j}) + g_4(T_{t-2,k_j}) + h(Toy_t) + \varepsilon_t,$$

où

- $Y_{t,j}^{det}$ est la charge électrique sans la tendance pour la zone j de l'observation t
- $DayType_t$ type de jours de l'observation t
- T_{t,k_j} la température de la station k_j associée à la zone j
- $\theta_{t,k_j} = (1 - 0.85)T_{t,k_j} + 0.85\theta_{t-1,k_j}$, T_{t-1,k_j} et T_{t-2,k_j} températures retardées respectivement de 24 heures et 48 heures.
- Toy_t position dans l'année de l'observation t

Cette méthodologie de modélisation a été performante et les auteurs de [91] ont fini la compétition troisième sur plus de 100 équipes.

Nous présentons ensuite le modèle initial utilisé sur les postes sources, qui est très proche du modèle moyen terme présenté précédemment.

Modèle pour les postes sources Par souci de lisibilité, nous ne faisons pas apparaître la dépendance en l'instant de la journée du modèle suivant, qui est issu de Goude *et al.* [55].

$$Y_{t,j} = \sum_{q=1}^9 m_q \mathbb{1}_{DayType_t=q} + \sum_{q=1}^{11} n_q \mathbb{1}_{Offset_t=q} + \mathbb{1}_{EJPSud_t=1} \\ + g_1(\theta_{t,k_j}) + g_2(T_{t,k_j}) + g_3(T_{t-1,k_j}) + g_4(T_{t-2,k_j}) \\ + h(Toy_t) + k(t) + \varepsilon_t,$$

où

- $Y_{t,j}$ charge consommée du poste source j pour l'observation t
- T_{t,k_j} la température de la station k_j associée au poste source j
- $\theta_{t,k_j} = (1 - 0.99)T_{t,k_j} + 0.99\theta_{t-1,k_j}$, T_{t-1,k_j} et T_{t-2,k_j} températures retardées respectivement de 24 heures et 48 heures.
- $Offset_t$: découpe l'année en plusieurs sous-périodes homogènes en termes de consommation
- $EJPSud_t$: Jour EJP ou non

Ce modèle est actuellement en phase de déploiement à ERDF.

2.2.3 Application des procédures de sélection automatique sur GEFCom 2012

Généralement, la charge consommée est fortement dépendante de la part climatique. Nous avons à disposition les séries de températures de plusieurs stations météorologiques. Ceci conduit aux deux axes d'étude suivants :

- Quels impacts a le choix d'une ou de plusieurs stations météorologiques ?
- L'automatisation sur toutes les zones de la sélection de stations météorologiques est-elle efficace ?

Nous appliquons des méthodes automatiques à cause du nombre important de séries chronologiques (une par heure et par zone). La correction court terme est importante car la variabilité plus grande des données conduit à des modèles moyen terme moins performants que pour le portefeuille EDF. Nous nous concentrons d'abord sur une zone.

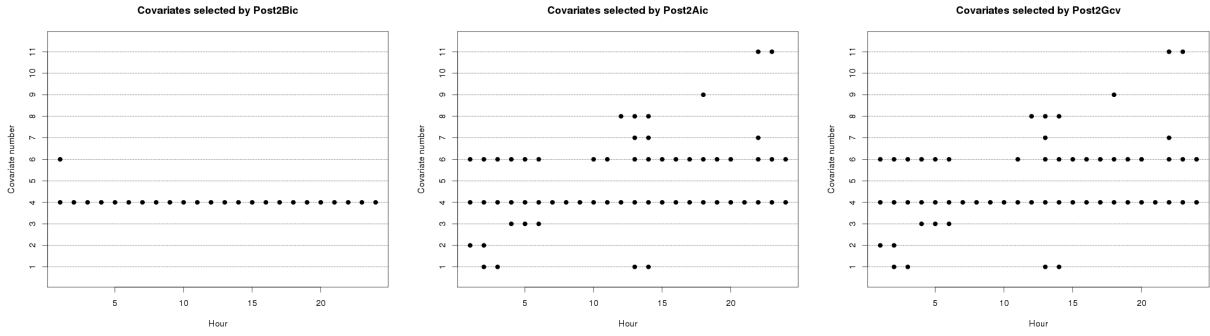


FIGURE 2.19 – Stations sélectionnées pour respectivement Post2Bic, Post2Aic et Post2Gcv selon l’instant sur la zone 14

2.2.3.1 Application à la zone 14

Benchmark Nous utilisons le modèle suivant, qui est inspiré du modèle moyen terme de [91], auquel une composante modélisant la tendance est ajouté et les températures retardées d’un et de deux jours supprimées. Comme pour les autres modèles, nous ne faisons pas apparaître la dépendance en l’heure de la journée du modèle.

$$Y_t = \sum_{q=1}^7 m_q \mathbb{1}_{DayType_t=q} + h(Toy_t) + k(t) + \sum_u^{11} l_u(\theta_{t,u}) + \sum_u^{11} g_u(T_{t,u}) + \varepsilon_t. \quad (2.5)$$

Ce modèle est appelé modèle saturé.

Phase de sélection Nous sélectionnons les stations météorologiques pour conserver la cohérence géographique, c’est-à-dire que soit les températures brute et lissée d’une station sont toutes les deux sélectionnées, soit aucune ne l’est.

La Figure 2.19, qui se lit comme la Figure 2.5, récapitule les stations météorologiques sélectionnées selon l’heure de la journée. Plusieurs stations météorologiques peuvent être sélectionnées, car les groupes de consommateurs reliés à la zone peuvent être éloignés les uns des autres. La sélection des stations d’un instant à l’autre est relativement homogène, ce qui est cohérent avec l’aspect local de l’application. Post2Bic sélectionne moins de stations que Post2Aic et Post2Gcv, ce qui illustre les propriétés du BIC, de l’AIC et du GCV.

Performance moyen terme Pour obtenir des modèles benchmarks, nous estimons tous les modèles comportant d’une à onze stations (qui est alors le modèle saturé).

$$Y_t = \sum_{q=1}^7 m_q \mathbb{1}_{DayType_t=q} + h(Toy_t) + k(t) + \sum_{i \in \mathbf{S}} l_i(T_{t,i}) + \sum_{i \in \mathbf{S}} g_i(\theta_{t,i}) + \varepsilon_t,$$

où $\mathbf{S}$ est l’ensemble des stations météorologiques candidates. Nous faisons cette recherche exhaustive, qui est longue et coûteuse informatiquement, pour comparer les différentes combinaisons de modèles qu’autorise le jeu de données. Nous vérifions que Post2 sélectionne automatiquement des combinaisons de stations conduisant à des modèles performants. La Figure 2.20 présente les boîtes à moustaches des MAPE en prévision des modèles obtenus avec la recherche exhaustive, ainsi que ceux de Post2Bic, Post2Aic et du Group LASSO sans correction (GrpL). Post2Gcv n’est pas représenté par souci de lisibilité, car ses performances sont très proches de celles de Post2Aic, ce qui est cohérent au vu de la proximité des modèles sélectionnés (voir la Figure 2.19).

Pour chaque nombre de stations, le point rouge représente le modèle sélectionné grâce au GCV au cours de la recherche exhaustive.

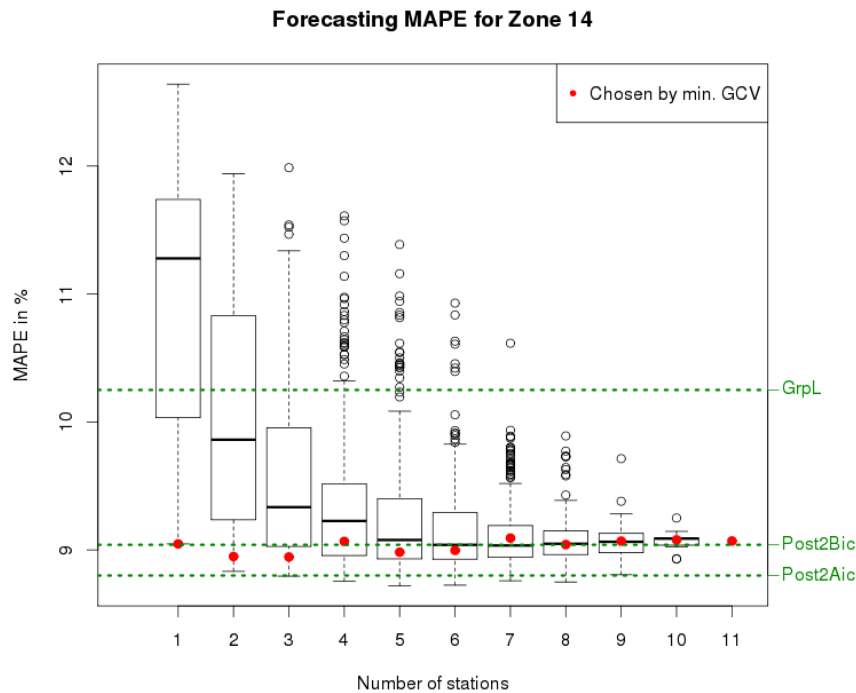


FIGURE 2.20 – MAPE selon les modèles considérés

La plus faible corrélation entre les températures des stations de la compétition est du niveau de celle entre les températures de, par exemple, Limoges et Clermont-Ferrand entre 2010 et 2012 alors que la plus forte est supérieure à celle entre les températures de, par exemple, Nîmes et Orange sur cette même période. Malgré ces fortes corrélations entre les températures, le choix d'une station ou d'une autre conduit à des qualités de prévision différentes, comme l'illustre la boîte à moustaches associée aux modèles comportant une station sur la Figure 2.20, car la température est un phénomène relativement localisé. En choisissant une température inadéquate, nous ne captons pas les épiphénomènes (par exemple, un orage, qui est un phénomène localisé, peut rapidement abaisser la température) qui influencent fortement la charge consommée. Une sélection efficace des stations est donc nécessaire. Moins il y a de stations sélectionnées, plus les performances en prévision risquent d'être mauvaises. Certaines stations sont éloignées des groupes de consommateurs. L'utilisation de plusieurs stations offre plus de flexibilité et intègre le fait que les groupes de consommateurs associés à une zone peuvent être éloignés. Malgré le problème de concurvité, qui est l'équivalent pour les modèles additifs des problèmes de multicollinéarité pour les modèles linéaires, les performances des modèles comportant beaucoup de stations météorologiques ne sont pas trop dégradées.

Les performances en prévision sont moins bonnes lorsque le biais du Group LASSO n'est pas corrigé (GrpL).

Post2Bic, Post2Aic et Post2Gcv ont des performances légèrement meilleures que celles du modèle saturé, ce qui illustre l'effet de sur-apprentissage. Post2Aic et Post2Gcv ont des meilleures performances que Post2Bic, illustrant de nouveau des propriétés propres à ces trois critères.

Avec onze stations, environ 5.5% du nombre d'estimateurs estimés pour la recherche exhaustive est nécessaire pour Post2. Malgré cela, Post2 a des performances prédictives équivalentes aux meilleurs modèles de la recherche exhaustive. Le modèle saturé a des performances satisfaisantes, mais perd en interprétabilité. Les procédures Post2 apportent plus d'informations, puisqu'elles

montrent que les stations 4 et 6 permettent d'avoir un modèle plus performant en termes de prévision que le modèle saturé et donc que, vraisemblablement, ces stations météorologiques sont proches des groupes de consommateurs de la zone 14.

Les études des RMSE et des MAPE en prévision, par mois et par heure, ainsi que les résidus justifient l'intérêt d'introduire des corrections court terme. Avril et mai sont les deux mois avec le MAPE le plus fort et où la consommation est la plus faible. Les températures de ces deux mois sont intermédiaires. Pour ces mois, la prévision pourrait être améliorée en intégrant des variables permettant de tenir compte de la plus grande variabilité de la température pendant l'inter-saison, en utilisant par exemple des écarts de températures. Les erreurs sont principalement commises autour de 10h et de 20h, ce qui correspond aux pics de consommation. En prévision, la part chauffage n'a pas été entièrement corrigée.

Nous étudions ensuite la qualité de l'automatisation de la sélection des stations météorologiques sur toutes les zones.

2.2.3.2 Application sur toutes les zones

Nous vérifions que nos procédures obtiennent des modèles performants sans intervention préalable ni hypothèse sur le nombre de stations.

Automatisation de la sélection des stations météorologiques pour le moyen terme

La recherche exhaustive étant trop longue pour être généralisée à l'ensemble des zones, nous introduisons trois méthodologies de sélection de stations météorologiques.

Pour la première méthode, nous utilisons le modèle saturé (2.5), c'est-à-dire que nous imposons un nombre fixe de stations et choisissons de toutes les utiliser, comme Lloyd [84]. Nous avons constaté les bonnes performances de cette méthode pour la zone 14.

Pour les deux autres méthodes, nous utilisons le modèle suivant (par souci de lisibilité, la dépendance en l'heure de la journée n'est pas indiquée) :

$$Y_{t,k} = \sum_{q=1}^7 m_q \mathbb{1}_{DayType_t=q} + h(Toy_t) + k(t) + g_1(\theta_{t,k}) + g_2(T_{t,k}) + \varepsilon_t.$$

La question des signaux météorologiques mis en entrée de ce modèle se pose. Hong *et al.* [63] proposent une sélection pas à pas de moyennes de températures. Cette méthode ne contraint pas le nombre de stations météorologiques et modélise l'effet de la température en utilisant en entrée un signal agrégé de températures. Elle peut être appliquée à plusieurs types de modèles. Nous nous inspirons aussi de Ben Taieb et Hyndman [17] et sélectionnons la station météorologique dont le signal de températures minimise le MAPE sur un échantillon de validation. Pour cette méthode, nous imposons la sélection d'une seule station météorologique et utilisons en entrée du modèle une température non agrégée.

Dans la Table 2.4, nous appelons ces trois méthodes respectivement Saturated, AggS et OneS. Lorsque Saturated est utilisé, les MAPE sont respectivement de 29.41% et de 65.20% pour les zones 9 et 10.

Zone	Saturated	Post2Bic	Post2Aic	Post2Gcv	AggS	OneS
Zone 1	7.62	8.52	7.90	7.89	8.49	8.49
Zone 2	3.79	3.86	3.77	3.81	3.90	4.04
Zone 3	3.79	3.86	3.77	3.81	3.90	4.04
Zone 4	41.44	41.06	41.03	41.04	41.48	41.28
Zone 5	6.50	6.76	6.66	6.64	6.91	6.91
Zone 6	3.68	3.75	3.66	3.70	3.81	3.99
Zone 8	6.69	6.96	6.79	6.81	6.95	6.95
Zone 11	6.37	7.15	6.91	6.91	7.25	7.25
Zone 12	4.40	4.54	4.37	4.37	4.45	4.72
Zone 13	7.03	6.88	6.78	6.78	6.88	6.88
Zone 14	9.07	9.04	8.80	8.82	8.98	9.05
Zone 15	8.01	7.73	7.74	7.74	7.80	7.80
Zone 16	7.44	7.64	7.33	7.37	7.69	8.32
Zone 17	5.18	5.60	5.24	5.24	5.24	5.46
Zone 18	5.54	5.75	5.56	5.57	6.00	6.53
Zone 19	8.28	8.29	8.24	8.25	8.43	8.44
Zone 20	4.42	4.44	4.37	4.37	4.40	4.50

TABLE 2.4 – MAPE ($\times 100$) pour les modèles moyen terme

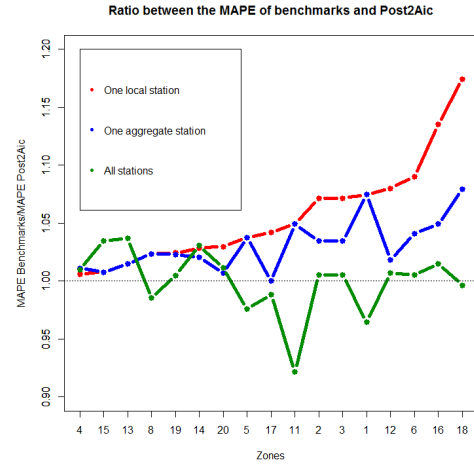


FIGURE 2.21 – Pour chaque zone, rapports entre les MAPE de OneS et celui de Post2Aic (rouge), celui d’AggS et celui de Post2Aic (bleu) ainsi que celui de Saturated et celui de Post2Aic (vert)

Post2Bic a généralement de moins bonnes performances en prévision que le modèle saturé, Post2Aic et Post2Gcv. Pour la majorité des zones, Post2Bic utilise moins de stations météorologiques, ce qui conduit à une moins grande flexibilité. Le modèle saturé, Post2Aic et Post2Gcv ont des performances très proches, bien que Post2Aic soit légèrement meilleur pour la majorité des zones. Les effets du modèle saturé peuvent être mal appris à cause de la forte corrélation entre les séries de températures. Utiliser la méthode de Hong *et al.* [65] présente l’avantage d’être rapide, simple et d’obtenir des bonnes performances. Cependant, utiliser une somme d’effets de températures et non un effet d’une seule série de températures (agrégée ou pas) est plus efficace, car cela conduit à plus de flexibilité. De plus, nous justifions théoriquement la consistance en sélection de nos familles de procédures dans le chapitre 3. Post2Aic sélectionne généralement moins de stations météorologiques que la méthode de [65] pour éviter des problèmes éventuels de sur-paramétrisation.

Apport de la correction court terme Les figures de cette section sont fournies dans l’Annexe C.4.2. Nous utilisons deux approches différentes pour la correction court terme de cette application. L’équation (2.6) donne le premier modèle de correction :

$$\hat{\varepsilon}_t = \beta_0 + \sum_{i=1}^7 \beta_i \hat{\varepsilon}_{t-i \times 24} + \mu_t + u_t, \quad (2.6)$$

où μ_t est un processus de moyenne mobile et u_t un bruit blanc. Nous apprenons d’abord le modèle sur l’échantillon d’apprentissage du moyen terme pour estimer $(\beta_i)_i$. La moyenne μ_t est alors considérée comme constante et est mise à jour séquentiellement durant le processus de prévision pour corriger les différents biais. La correction court terme du modèle saturé améliore les performances pour toutes les zones (excepté la zone 4, qui est spéciale), le rapport entre les MAPE des prévisions court et moyen termes étant entre 0.55 et 0.9. La correction court terme réduit la variance des erreurs et centre les erreurs. Comme nous travaillons sur des charges non corrigées de la tendance et plus variables, l’utilisation d’un processus de moyenne mobile améliore les performances de prévision.

Nous nous intéressons aussi à l'évolution de la comparaison des performances entre le moyen et le court terme. Le court terme réduit les écarts de performance entre les différents modèles moyen terme. Cependant, la comparaison des performances de certaines zones change par rapport au moyen terme. Les zones 1, 5, 11 et 17, qui sont mal prédites par Post2Aic par rapport au modèle saturé à moyen terme, sont mieux prévues après correction court terme. Ces trois dernières zones font partie des quatre zones les mieux corrigées lors du passage au court terme pour le modèle saturé. Le biais à moyen terme est élevé (la tendance peut avoir été mal corrigée). Lorsque nous corrigeons artificiellement ce biais, Post2Aic a de meilleures performances moyen terme que le modèle saturé. Le court terme corrige une partie du biais, ce qui explique les meilleures performances de la version court terme de Post2Aic par rapport à celles du modèle saturé. Post2Aic court terme a de meilleures prévisions pour la majorité des zones.

La seconde approche de correction consiste en l'application de la procédure Post1 (adaptée au modèle linéaire) pour sélectionner un modèle auto-régressif pour les erreurs du modèle moyen terme. Ceci consiste schématiquement à sélectionner les erreurs retardées qui servent de covariables pour le modèle de correction court terme. Nous utilisons le modèle saturé présenté dans l'équation (2.7) :

$$\begin{aligned} \hat{\varepsilon}_t = & \beta_0 + \sum_{i=0}^{143} \beta_{i+1} \hat{\varepsilon}_{t-24-i} + \sum_{i=1}^6 \alpha_i \left(1/24 \sum_{j=0}^{23} \hat{\varepsilon}_{t-24 \times i - j} \right) + \sum_{i=1}^6 \gamma_i \max_{j \in \llbracket 0, 23 \rrbracket} (\hat{\varepsilon}_{t-24 \times i - j}) \\ & + \sum_{i=1}^6 \delta_i \min_{j \in \llbracket 0, 23 \rrbracket} (\hat{\varepsilon}_{t-24 \times i - j}) + \sum_{i=1}^6 \theta_i \text{median}_{j \in \llbracket 0, 23 \rrbracket} (\hat{\varepsilon}_{t-24 \times i - j}) + \phi_1 \left(\sum_{i=1}^6 \alpha_i (1/24 \sum_{j=0}^{23} \hat{\varepsilon}_{t-24 \times i - j}) \right) \\ & + \phi_2 \max_{j \in \llbracket 0, 143 \rrbracket} (\hat{\varepsilon}_{t-24-j}) + \phi_3 \min_{j \in \llbracket 0, 143 \rrbracket} (\hat{\varepsilon}_{t-24-j}) + \phi_4 \text{median}_{j \in \llbracket 0, 143 \rrbracket} (\hat{\varepsilon}_{t-24-j}) + \mu_t + u_t. \end{aligned} \quad (2.7)$$

Ce modèle contient des erreurs retardées brutes (classique dans le cadre de la correction court terme) et des erreurs retardées journalières et hebdomadaires (maximum, minimum, ...) qui servent à diminuer la variabilité des erreurs. μ_t et u_t sont respectivement un processus de moyenne mobile et un bruit blanc.

Pour la correction de Post2Aic, les erreurs passées les plus souvent sélectionnées sont celles de précisément 1, 2, 6 et 7 jours, ainsi que celles retardées de 25 heures (voir Annexe C.4.2), ce qui est cohérent, car pour ces retards, les niveaux de charges et des erreurs correspondent mieux. Les erreurs accessibles médianes de la semaine et de la veille sont sélectionnées pour la majorité des zones. Elles présentent l'avantage d'être moins variables que les erreurs passées, apportent une information sur la complexité des prévisions récentes et corrigent mieux les biais éventuels. Les erreurs maximum et minimum journalières et hebdomadaires sont régulièrement sélectionnées. Elles apportent une information sur l'intensité des erreurs récentes. Les tailles des modèles sélectionnés sont disparates. La correction des modèles les mieux prédits à moyen terme a tendance à être moins parcimonieuse, ce que nous expliquons par un apport moindre d'informations par les erreurs pour ces modèles.

La sélection des retards et l'ajout d'informations sur les erreurs médianes, maximales et minimales journalières et hebdomadaires, qui sont plus lisses que les erreurs instantanées du passé, améliorent la prévision pour toutes les zones, avec une diminution des MAPE de 1 à 6%. Ici, le BIC est plus efficace que le GCV ou l'AIC pour sélectionner les retards, car la présence de beaucoup de retards conduit à un effet de sur-apprentissage. Les erreurs de deux instants voisins peuvent être fortement corrélées, posant des problèmes d'identifiabilité.

Fort de ces différentes observations, nous passons à l'étude d'une maille locale d'un autre réseau électrique : les postes sources.

2.2.4 Application sur les postes sources

Dans cette sous-section, nous étudions les postes sources de l'ACR de Lyon. Nous avons exclu de l'étude une quinzaine de postes sources au comportement erratique, avec des valeurs isolées proches de zéro ou avec une forte rupture de la consommation et travaillons avec 61 postes sources. Nous invalidons toutes les données dont la charge consommée est nulle. Par souci de lisibilité, nous ne faisons pas apparaître la dépendance en l'instant dans les modèles de cette section. Pour le poste source j , le modèle Benchmark est donné par la formule suivante :

$$\begin{aligned} Y_{t,j} = & \sum_{q=1}^9 m_q \mathbb{1}_{DayType_t=q} + \sum_{q=1}^{11} m_q \mathbb{1}_{Offset_t=q} + \mathbb{1}_{EJPSud_t=1} \\ & + g_1(\theta_t) + g_2(T_t) \\ & + h(Toy_t) + k(t) + \varepsilon_t, \end{aligned}$$

où les températures brutes T_t et lissées θ_t sont mesurées à Lyon.

Nous considérons Mulhouse, Dijon, Nevers, Limoges, Clermont-Ferrand, Lyon, Bourg-Saint-Maurice, Nîmes, Orange, Le Luc et Saint-Auban comme les stations météorologiques candidates. Ce sont les stations les plus proches de l'ACR Lyon dont nous avons les mesures et elles apportent des informations aux quatre points cardinaux de l'ACR. Nous appliquons la procédure automatique Post2Aic pour réaliser le même type d'études que pour la compétition GEFCom 2012, c'est-à-dire que nous cherchons à sélectionner les stations météorologiques. Comme pour l'application sur les données GEFCom 2012, nous groupons les températures lissées et brutes par station météorologique. Nous nous concentrons sur Post2Aic, parce que cette méthode a montré sur les deux précédents jeux de données des performances légèrement meilleures que Post2Bic et Post2Gcv. Pour chaque heure de la journée (la dépendance dans l'instant n'apparaît pas ici) et chaque poste source j , Post2Aic va sélectionner un modèle du type :

$$\begin{aligned} Y_{t,j} = & \sum_{q=1}^9 m_q \mathbb{1}_{DayType_t=q} + \sum_{q=1}^{11} n_q \mathbb{1}_{Offset_t=q} + \mathbb{1}_{EJPSud_t=1} \\ & + \sum_{b \in \mathbf{S}_j} g_b(\theta_{t,b}) + \sum_{b \in \mathbf{S}_j} l_b(T_{t,b}) \\ & + h(Toy_t) + k(t) + \varepsilon_t, \end{aligned}$$

où $\mathbf{S}_j$ est l'ensemble des stations météorologiques sélectionnées par Post2Aic pour le poste source j (à l'instant considéré).

La Figure 2.22 illustre l'application, l'éclair représentant le placement (approximatif) de quelques postes sources de l'ACR et le soleil avec le nuage les stations météorologiques. Dans cette application, nous nous intéressons à la capacité des procédures à sélectionner automatiquement des modèles performants et répondons à la question de l'apport de la sélection des points de mesure de la température par rapport à l'utilisation d'une station prédéterminée (pour cette ACR, celle de Lyon). La Figure 2.23 présente les rapports entre les MAPE en prévision de Post2Aic et les MAPE en prévision du modèle Benchmark, ainsi que les boîtes à moustaches qui leur sont associées. Si ces rapports sont inférieurs à 1, alors les procédures de sélection de stations météorologiques améliorent les modèles. 3% de ces rapports sont supérieurs à 1.05. Pour ces deux postes sources, qui sont situés à Lyon ou en proche banlieue, l'utilisation de la température mesurée à Lyon conduit à des performances fortement meilleures. Inversement, deux tiers des rapports sont inférieurs à 1 et 23% sont inférieurs à 0.95. Pour ces derniers postes sources, l'application de la procédure Post2Aic a donc fortement amélioré la prévision. Tous ces postes sources utilisent pour



FIGURE 2.22 – Illustration de l'application sur les postes sources

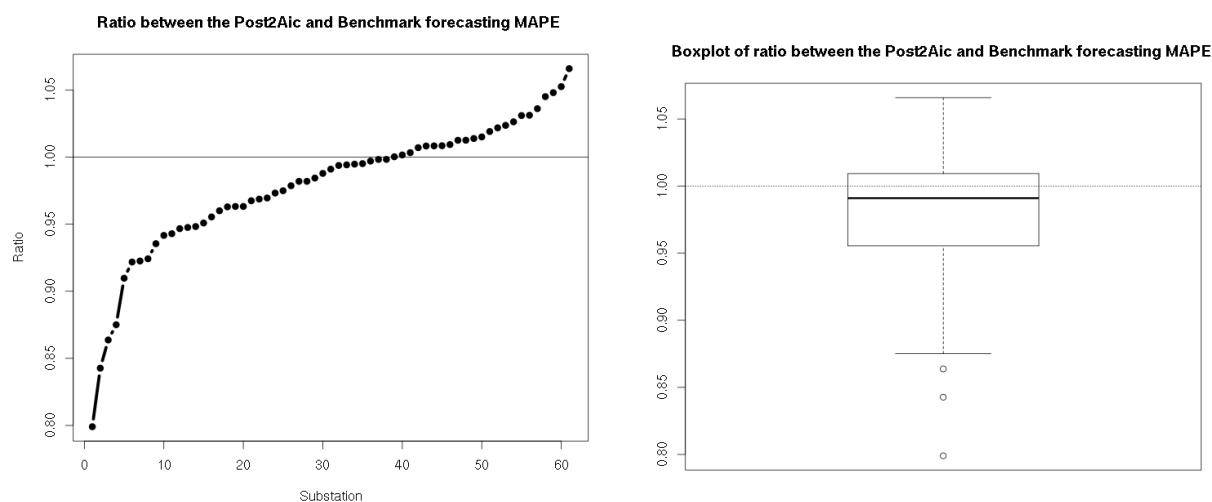


FIGURE 2.23 – Rapport entre les MAPE en prévision de Post2Aic et les MAPE en prévision du modèle Benchmark

la grande majorité des heures les températures de Lyon et Clermont-Ferrand, qui sont les deux grandes villes à proximité de l'ACR Lyon dont nous avons la météorologie. Les postes sources de l'Ain mieux prédits par Post2Aic par rapport au Benchmark sélectionnent régulièrement des stations telles que Nevers, Dijon ou Mulhouse, alors que ceux de l'Isère sélectionnent régulièrement Bourg-Saint-Maurice. Pour les autres postes sources dont les rapports des MAPE sont inférieurs à 0.95, les stations météorologiques des grandes villes sont régulièrement sélectionnées, alors que les stations météorologiques situées à Orange, Le Luc ou Saint-Auban sont peu sélectionnées. Un poste source peut être influencé par un poste source voisin et donc indirectement par des événements plus ou moins lointains. Deux postes sources situés à Lyon ou en proche banlieue font partie des postes sources dont le rapport des MAPE est inférieur à 0.95. Ces postes sources ont une connexion de tension inhabituelle par rapport à la majorité des postes sources.

L'utilisation de plusieurs stations météorologiques améliore les prévisions de consommation sur certains postes sources de l'ACR Lyon. Cependant, nous restons limités par le fait que nous n'avons pas à disposition pour cette étude de stations météorologiques plus proches et plus nombreuses dans la zone (Saint-Étienne, Grenoble ou Bourg-en-Bresse).

Conclusion

Nous faisons ici une rapide synthèse des résultats et présentons des pistes d'amélioration pour les deux types d'études.

Synthèse des résultats de l'étude du portefeuille EDF Pour cette application, nous comparons les modèles sélectionnés par nos procédures avec un modèle additif de prévision moyen terme de la consommation électrique qui a été optimisé par un expert d'EDF R&D (dit modèle EDF). Pour obtenir son modèle, l'expert fait appel à son expertise métier et teste un nombre important de combinaisons de modèles. Même sur un jeu de données similaire, l'expert aurait à refaire toute l'optimisation. Le modèle EDF est performant mais son optimisation est longue, fastidieuse et difficilement généralisable. L'enjeu de l'étude est d'obtenir automatiquement des performances équivalentes. Nous appliquons les procédures Post2 (en utilisant le BIC, l'AIC et le GCV) pour automatiser la procédure de sélection. Nous nous intéressons plus particulièrement à la sélection des variables météorologiques. Nous partageons notre échantillon de données en deux échantillons : un échantillon de quelques années où les modèles sont estimés et un échantillon de quelques mois où nous testons les performances des modèles. Nous travaillons à l'horizon moyen terme et vérifions que les comparaisons pour cet horizon restent valables à court terme (avec un horizon de 24 heures). Nous estimons un modèle différent par demi-heure, car les effets des covariables sont différents selon l'instant, ce qui conduit à supposer que les instants sont indépendants. Nous corrigeons au préalable la tendance en estimant le modèle EDF sur l'échantillon total (apprentissage+test) puis en retranchant à la charge les effets estimés de la tendance.

Comme nous travaillons instant par instant, les sous-ensembles de covariables sélectionnées par Post2Bic, Post2Aic et Post2Gcv sont différents selon la demi-heure considérée. Post2Bic sélectionne moins de covariables que Post2Aic et Post2Gcv, qui sont proches. Ceci est cohérent avec les propriétés connues du BIC, de l'AIC et du GCV. Il y a des similitudes et des différences entre les covariables sélectionnées par l'expert et les procédures issues de Post2. Ces dernières sélectionnent plus de covariables. Il y a une explication physique aux covariables sélectionnées. Par exemple, la nébulosité est principalement sélectionnée le jour et non la nuit, ce qui est cohérent, car la nébulosité influence fortement l'éclairage. Les sous-ensembles de covariables sélectionnées sont relativement homogènes d'un instant à l'autre. Nous utilisons la dépendance entre les instants pour post-traiter les sous-ensembles de covariables sélectionnés. L'un des post-traitements consiste à conserver un sous-ensemble de covariables commun à tous les instants. Pour cela, nous initialisons

une grille de pourcentages de présence d'une covariable, estimons les modèles associés à chaque pourcentage et sélectionnons le modèle associé grâce à une technique du coude sur le RMSE en estimation.

Nous comparons les performances du modèle EDF et celles des modèles issus de la procédure Post2 avec des modèles additifs génériques pour montrer l'importance de la sélection de variables pour améliorer les performances. Le modèle EDF a des performances équivalentes voire légèrement moins bonnes que celles des modèles issus de Post2. Le post-traitement de la sélection apporte des forts gains, notamment pour Post2Bic : la version non traitée de Post2Bic est la moins performante alors que sa version post-traitée conduit au modèle le plus performant. Lorsque nous étudions les erreurs en fonction du mois, nous constatons que les modèles prévoient moins efficacement deux mois. Pour un de ces mois (novembre), il y a une rupture de la consommation électrique et la charge consommée est anormalement basse. Pour l'autre (avril), les températures semblent plus froides que la normale, expliquant en partie ces erreurs. Les erreurs sont corrélées entre elles. Nous utilisons cette information pour corriger les modèles moyen terme à un horizon court terme (charges retardées de 24 heures supposées connues). Après correction, les performances sont très bonnes avec des MAPE en prévision inférieurs à 1%. Les améliorations sont dues à la correction de biais et à la réduction de la variance. Même après ce passage au court terme, les comparaisons de performances restent similaires. Bien que l'amélioration soit faible, les modèles issus de Post2 permettent d'avoir des meilleures performances que le modèle EDF. Les post-traitements de Post2 sont utiles, même après le passage au court terme. Le mois de novembre est bien corrigé à court terme, contrairement à avril. Ceci est cohérent avec ce qui nous semble être l'origine des erreurs pour ces deux mois.

Les gains de performance apportés par Post2 sont faibles, mais la procédure présente le fort avantage d'être automatique et de demander beaucoup moins d'interventions humaines et de connaissances métiers sur le portefeuille EDF.

Perspectives de l'étude du portefeuille EDF Nous avons travaillé avec des modèles par instant. Adapter nos procédures aux modèles à coefficients variables selon l'instant de la journée permet de s'affranchir de cette hypothèse et de travailler avec des modèles globaux. Des hypothèses devraient alors être faites sur la matrice de covariance.

Nous avons établi des procédures automatiques de sélection de variables pour des modèles additifs que nous corrigeons ensuite à court terme. Une piste d'amélioration est de rendre les méthodes adaptatives à court terme. Ceci pose cependant deux questions :

- Les modèles additifs peuvent-ils facilement être adaptatifs ?
- Comment rendre les procédures de sélection adaptatives ?

Une autre piste d'amélioration est de travailler sur les signaux de températures utilisés. Pour cette étude, nous travaillons sur une température nationale qui consiste en une moyenne pondérée électriquement de 32 stations météorologiques. Au lieu de cette température nationale, nous pourrions utiliser les températures de chaque station météorologique directement en entrée des modèles.

Synthèse de l'étude des données locales Nous travaillons sur deux types de données de prévision de consommation électrique locale. Les premières sont issues de la compétition GEFCom 2012 proposée par le site Kaggle et les secondes sont propres à un niveau local du réseau électrique français, les postes sources. Les études locales imposent d'avoir des méthodes automatiques, car il y a de nombreuses séries temporelles à rapidement étudier. La question de la localisation des mesures des covariables météorologiques est importante, mais peu étudiée. Souvent, pour étudier un niveau local de consommation électrique, les auteurs l'abordant proposent un nombre fixé de stations météorologiques. Nous nous affranchissons de cette hypothèse et utilisons nos méthodes de sélection automatique de variables dans les modèles additifs. Nous avons montré

empiriquement que la sélection des stations météorologiques est importante pour bien prévoir les charges consommées, malgré la forte corrélation entre les températures, ce qui nous semble cohérent, car il y a des épiphénomènes localisés pour la température qui influencent fortement la charge consommée. S'affranchir d'un nombre fixé de stations à utiliser améliore les prévisions, car cela offre une plus grande flexibilité permettant de tenir compte de la structure du réseau. À cause de ce besoin de flexibilité, la qualité de prévision est améliorée par l'utilisation de plusieurs signaux de températures, et non d'un unique, même si celui-ci est une agrégation de plusieurs signaux de températures. Post2 répond à la fois à ces nécessités et au besoin d'automatisation induit par l'étude d'une maille locale de consommation électrique. Pour ces applications, la correction court terme est importante. L'application de Post1 adaptée aux modèles linéaires améliore le correctif.

Perspectives de l'étude des données locales Les différents membres d'une maille d'un réseau interagissent. À court terme, utiliser l'information des charges consommées passées des autres zones ou postes sources peut améliorer les performances des modèles. Nous pouvons chercher à sélectionner les charges passées des postes sources ou zones voisins avec les procédures Post2 pour améliorer la prévision court terme. Concernant les postes sources, au cours d'une phase d'expérimentation préliminaire, où nous avons à disposition une grille plus fine des stations météorologiques (avec les températures de Saint-Étienne, Grenoble ou Mâcon par exemple), les procédures obtenaient des meilleures performances. Les variables explicatives candidates sont très corrélées. Un changement de pénalisation, qui tiendrait compte de cette corrélation, pourrait améliorer les performances. Dans le contexte linéaire, les auteurs de [43] combinent la pénalisation L1 avec une pénalisation tenant compte de la corrélation. Pour ces données, nous avons travaillé sur la sélection des points de la maille du réseau où sont mesurées les variables de températures. Nous aurions pu créer à partir des températures mesurées à différents endroits un dictionnaire de covariables et faire le même type d'études que pour le portefeuille EDF. Pour avoir une vision probabiliste, nous aurions pu remplacer le coût quadratique du Group LASSO et des P -Splines par la perte quantile [74].

Un résultat de consistance pour un estimateur en plusieurs étapes

3.1 Introduction

Les modèles additifs, introduits par Hastie et Tibshirani [57], permettent de réaliser un bon compromis entre flexibilité et complexité. Nous utilisons des bases de B-Splines (Stone, [106]) pour en approximer ces composantes. Le problème de sélection de composantes est un sujet classique en statistique. Pour y répondre dans le cas des modèles additifs, Fan et Jiang [44] utilisent des tests de déviance. Marra et Wood [86] proposent deux versions modifiées des P-Splines [42]. Pour avoir des méthodes efficaces à la fois en termes de sélection et d'estimation, de nombreux auteurs utilisent des estimateurs en plusieurs étapes. Par exemple, Antoniadis *et al.* [9] utilisent un estimateur Nonnegative Garrote [20] ayant pour estimateur initial un estimateur P-Splines ou Huang *et al.* [67] utilisent un estimateur Group LASSO adaptatif, dont l'estimateur initial est un estimateur Group LASSO [122].

Nous nous intéressons (asymptotiquement) à la sélection et à l'estimation de composantes dans le modèle (3.1) :

$$Y_i = \beta_0^* + \sum_{j=1}^d f_j^*(X_{i,j}) + \epsilon_i, \quad (3.1)$$

avec $i = 1, \dots, n$, $E(\epsilon_i | \underline{X}_i) = 0$, $\underline{X}_i := (X_{i,1}, \dots, X_{i,d})^T$, ϵ_i *i.i.d.* de variance σ^2 inconnue et des fonctions f_j^* suffisamment régulières en un sens qui sera précisé ci-dessous. Ces hypothèses forment un cadre classique pour la régression additive non-paramétrique. Nous n'imposons pas de contraintes sévères sur la dimension d mais nous supposons qu'il n'existe qu'un certain nombre (inconnu) fini et borné de composantes f_j^* qui ont un effet significatif sur la régression. Le modèle (3.1) est donc un modèle additif creux. Nous supposons sans perte de généralité que le vecteur aléatoire $\underline{X} = (X_1, \dots, X_d)^T$ est de loi à support dans $[0, 1]^d$.

Nous supposons que les fonctions f_j^* appartiennent à un espace de Sobolev :

$$H_2^2([0, 1]) = \{f : [0, 1] \rightarrow \mathbb{R} \mid f^{(1)} \text{ absolument continue et } \int_0^1 (f^{(2)}(x))^2 dx < +\infty\}.$$

De plus, pour des raisons d'identifiabilité, nous supposons que pour tout j , $E(f_j^*(X_j)) = 0$. Nous notons

$$S^* = \{j \in \{1, \dots, d\} \mid E(f_j^*(X_j)^2) \neq 0\},$$

l'ensemble des composantes non nulles du modèle et $s^* = \text{card}(S^*)$ le cardinal de S^* . Nous supposons de plus que $s^* \ll n$ et $s^* \leq D < d = d(n)$. Avec ces hypothèses, nous supposons qu'il y a un nombre restreint de variables explicatives significatives comparé au nombre d'observations. Par contre, nous autorisons le nombre de variables candidates (la dimension du modèle) à augmenter en fonction du nombre d'observations.

Notre objectif est d'obtenir un estimateur qui soit creux, c'est-à-dire qu'il ne doit pas contenir trop de composantes additives non informatives, et qui ait un bon comportement asymptotique, c'est-à-dire que son erreur quadratique moyenne tende vers 0.

Nous avons choisi d'approcher les composantes fonctionnelles du modèle additif (3.1) par leur développement tronqué dans des bases de B-Splines [106]. Avec cette approximation, nous pouvons approcher le modèle (3.1) par le modèle (3.2) :

$$E(Y|X_1, \dots, X_d = (x_1, \dots, x_d)) = \beta_0 + \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(x_j), \quad (3.2)$$

avec $\mathbf{B}_j(\cdot) = (B_{j,k}^{q_j}(\cdot) | k = 1, \dots, K_j + q_j = m_j)$ la base de l'espace vectoriel engendré par les B-Splines de degré q_j et K_j noeuds, sur lequel est projetée la j ème composante additive $f_j^*(X_j)$. Soit $\mathbf{B}_j = (\mathbf{B}_j(x_{1j})^T, \dots, \mathbf{B}_j(x_{nj})^T)^T \in \mathbb{R}^{n \times m_j}$, où $x_{i,j}$ constitue la i ème observation de la variable X_j . Les paramètres à estimer du modèle sont alors β_0 et $\boldsymbol{\beta} = (\beta_{1,1}, \dots, \beta_{d,m_d})^T$. Les variables explicatives sont les $(X_i)_{i \in \llbracket 1, d \rrbracket}$. Elles peuvent être influentes, ou non, sur la réponse Y . Les composantes du modèle sont les f_j^* et les paramètres à estimer du modèle sont β_0 et $\boldsymbol{\beta} = (\beta_{1,1}, \dots, \beta_{d,m_d})^T$. Nous supposons que les m_j , avec $j \in \llbracket 1, d \rrbracket$, sont connus. Sauf mention contraire, nous admettrons que $m_1 = \dots = m_d = m_n = m$.

En projetant les composantes additives dans des bases de B-Splines, nous nous ramenons au cas d'un modèle linéaire. Lorsque nous retenons une covariable, nous devons donc sélectionner tout le groupe de coefficients qui lui est associé dans la base correspondante de B-Splines. Nous avons donc besoin d'une méthode qui sélectionne groupe par groupe. Nous utilisons le Group LASSO pour ajuster et sélectionner le nombre approprié de composantes additives du modèle.

Si le Group LASSO [122] permet de réaliser la sélection de composantes dans les modèles additifs, il a les mêmes inconvénients que le LASSO [112]. Du fait que la pénalisation sur les groupes de coefficients soit convexe (pénalité L1), ce qui est équivalent à un seuillage doux, les coefficients estimés, même les plus forts, sont sur-rétrécis, conduisant à un biais important sur l'estimation des composantes. Le fait que le degré de pénalisation soit constant quelque soit l'amplitude des coefficients explique en partie ce biais. Pour compenser cet effet, le Group LASSO a tendance à inclure des variables non pertinentes. Pour corriger ce biais, et pour mettre au point des méthodes simultanément optimales en sélection et en prédiction asymptotique, nous pouvons utiliser des estimateurs en plusieurs étapes. Nous avons trouvé peu de résultats théoriques démontrés concernant les estimateurs en plusieurs étapes. Antoniadis *et al.* [8] appliquent des variantes de nos procédures avec des bons résultats pratiques. Ceci motive notre étude.

Dans la suite, nous définissons les procédures d'estimateurs en plusieurs étapes que nous avons étudiées. $S = \{1, \dots, d\}$ contient les numéros associés aux variables explicatives présentes dans le dictionnaire de covariables. Supposons qu'il existe $\lambda_{\max}$ tel que tous les coefficients estimés par l'estimateur Group LASSO associé à ce λ soient nuls et $\lambda_{\min}$ tel que $\lambda_{\min} << \lambda_{\max}$. Nous construisons une grille de paramètres de lissage uniforme Λ_{GrpL} entre $\lambda_{\min}$ et $\lambda_{\max}$.

Les procédures de type Post, que nous étudions plus en détail par la suite, sont de la forme suivante :

Algorithme Post

1. **Première étape : Construction de sous-plans d'expérience candidats :** Pour chaque $\lambda_i \in \Lambda_{GrpL}$
 - Notons S_{λ_i} le sous-ensemble des numéros des covariables sélectionnées par l'estimateur associé au Group LASSO de paramètre λ_i .
 - Résoudre $\tilde{\beta}_{\lambda_i} = \arg \min \sum_{l=1}^n \left(Y_l - \beta_0 - \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(X_{l,j}) \right)^2 + \lambda_i \sum_{j=1}^d \sqrt{m_j} \|\beta_j\|_2$.
 - Pour chaque $j \in S$, si $\|\tilde{\beta}_{\lambda_i,j}\|_2 \neq 0$, ajouter j dans S_{λ_i} , sinon ne pas considérer j dans S_{λ_i} . Pour faciliter les notations, nous supposons que tous les sous-plans d'expérience candidats sont différents.
2. **Seconde étape : Estimation des modèles candidats :** Pour chaque $S_{\lambda_s} \in \{S_{\lambda_{min}}, \dots, S_{\lambda_{max}}\}$
 - Résoudre un problème du type $\hat{\beta}_{S_{\lambda_s}} = \arg \min Q$ où Q est la fonction objective des moindres carrés ordinaires (MCO) ou des P-Splines [42].
 - Calcul du critère de sélection de modèle (CSM), typiquement le BIC [102], l'AIC [1] ou le GCV [115], de l'estimateur $\hat{\beta}_{S_{\lambda_s}}$.
3. **Troisième étape : Sélection de l'estimateur final :** Sélectionner $\hat{\beta}_{S_{\lambda_b}}$ qui minimise le CSM choisi

Les procédures de type Ante sont de la forme suivante :

Algorithme Ante

1. **Première étape : Construction de sous-plans d'expérience candidats :** Pour chaque $\lambda_i \in \Lambda_{GrpL}$
 - Notons S_{λ_i} le sous-ensemble des indices des covariables sélectionnées par l'estimateur associé au Group LASSO de paramètre λ_i .
 - Résoudre $\tilde{\beta}_{\lambda_i} = \arg \min \sum_{l=1}^n \left(Y_l - \beta_0 - \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(X_{l,j}) \right)^2 + \lambda_i \sum_{j=1}^d \sqrt{m_j} \|\beta_j\|_2$
 - Pour chaque $j \in S$, si $\|\tilde{\beta}_{\lambda_i,j}\|_2 \neq 0$, ajouter j dans S_{λ_i} , sinon ne pas considérer j dans S_{λ_i} .
 - Calcul du critère de sélection de modèle (CSM), typiquement le BIC, l'AIC ou le GCV, de l'estimateur $\tilde{\beta}_{\lambda_i}$.
2. **Deuxième étape : Sélection du plan d'expérience candidat final :** Sélectionner $\tilde{\beta}_{\lambda_i}$ qui minimise le CSM choisi, notons S_{λ_s} le plan d'expérience associé à l'estimateur.
3. **Troisième étape : Estimation du modèle sélectionné :** Résoudre un problème du type $\hat{\beta}_{S_{\lambda_s}} = \arg \min Q$ où Q est la fonction objective des MCO ou des P-Splines.

Dans le cadre usuel de la théorie de la statistique inférentielle, nous nous intéressons généralement aux propriétés asymptotiques des estimateurs sous les hypothèses d'un modèle donné a priori. Nous ne sommes pas dans ce cadre, puisque nos procédures comportent une étape de sélection de modèles fondée sur les données. Pour avoir des résultats à la fois en termes de sélection et d'estimation, il faut que les composantes non nulles soient uniformément plus grandes qu'une borne inférieure de rapport de signal sur bruit. Ceci n'est généralement pas vérifié en pratique. Si cette condition n'est pas vérifiée, une estimation consistante des composantes après sélection est généralement impossible. Leeb et Potscher [81] considèrent le problème qui consiste à estimer la distribution d'un estimateur conditionnellement au fait qu'on ait sélectionné au préalable un modèle. Ils proposent des procédures de type Ante (sélection d'un modèle, suivi d'une ré-estimation des paramètres) et montrent que l'estimateur ne peut généralement pas être consistant uniformément. L'hypothèse que les composantes non nulles soient plus fortes qu'un certain

niveau dépendant du bruit est intuitive : si certaines composantes sont trop proches de l'application nulle, elles seront plus difficiles à sélectionner et à estimer. Néanmoins, pour démontrer nos résultats, nous supposons seulement que la norme Euclidienne des composantes non nulles ne tend pas trop vite vers 0 asymptotiquement. Nous nous affranchissons donc de l'hypothèse classique qui consiste à supposer que la norme Euclidienne de chaque composante non nulle est supérieure à une constante strictement positive et indépendante de n .

Dans le chapitre 1, nous avons étudié empiriquement les propriétés des procédures utilisant le BIC, le GCV et l'AIC. Par la suite, nous nous concentrons sur l'application du BIC, pour lequel nous montrons des résultats asymptotiques. Notre objectif est d'obtenir un estimateur qui soit creux (dans le sens qu'il ne contient pas trop de composantes additives non informatives) et qui ait un bon comportement asymptotique. Ce problème est plus difficile à résoudre dans le cas non-paramétrique que dans le cas paramétrique, parce que la classe de fonctions dans laquelle nous cherchons à estimer la fonction de régression est plus compliquée. Il faut tenir compte dans l'étude théorique du fait que les coefficients que nous estimons permettent d'estimer une approximation des composantes des modèles additifs, et non les composantes elles-mêmes. Meier *et al.* [89] proposent, dans un article de référence, une méthode d'estimation par moindres carrés pénalisés en introduisant la pénalité (3.3) suivante :

$$\lambda_1 \sum_{j=1}^d \sqrt{\|f_j^*\|_{2n}^2 + \lambda_2 I(f_j^*)^2}, \quad (3.3)$$

où $\|f_j\|_{2n}^2 = 1/n \sum_{i=1}^n f_j(X_{i,j})^2$ et $I(f_j^*)^2 = \int_0^1 f_j^{*(2)}(x)^2 dx$. Ces deux termes pénalisent respectivement l'entrée de la covariable j et la non régularité de la composante f_j^* . Ils assurent le bon comportement asymptotique de l'estimateur et de l'absence de sur-apprentissage. Posons Ω_j tel que $\Omega_{j,kl} = \int B_{j,k}^{q_j(2)}(x) B_{j,l}^{q_j(2)}(x) dx$, $M_j = B_j^T B_j / n + \lambda_2 \Omega_j$. A l'aide d'une décomposition de Cholesky, il existe R_j telle que nous puissions écrire $M_j = R_j^T R_j$. Nous notons $\tilde{\beta}_j = R_j \beta_j$ et $\tilde{\mathbf{B}}_j = \mathbf{B}_j R_j^{-1}$. Notons $\tilde{\mathbf{B}} = (\tilde{\mathbf{B}}_1 | \dots | \tilde{\mathbf{B}}_d)$. Meier *et al.* [89] montrent que le problème d'optimisation des moindres carrés pénalisés par la pénalisation (3.3) est alors équivalent à optimiser :

$$\arg \min \|\mathbf{Y} - \tilde{\mathbf{B}} \tilde{\beta}\|_{2,n}^2 + \lambda_1 \sum_{j=1}^d \sqrt{m_j} \|\tilde{\beta}_j\|_2,$$

et nous reconnaissons là un problème de type Group LASSO. Meier *et al.* [89] proposent de sélectionner λ_2 dans une grille de taille 15. Nous avons remarqué sur une simulation, présentée dans l'Annexe D, qu'il n'y a pas de faux négatifs, mais qu'il y a des faux positifs avec cette méthode. En plus de la possibilité d'introduire des faux positifs, en pratique, cette double pénalisation conduit à des résultats peu satisfaisants en termes de prévision car l'estimateur résultant est moins performant que l'estimateur oracle, qui serait un estimateur construit lorsque S^* est supposé connu, et ceci vaut même si le vrai modèle est sélectionné. Cela semble justifier le fait de séparer la phase de sélection et la phase d'estimation.

Dans le cas des modèles de régression linéaire multiple, Belloni *et al.* [16] ont considéré une procédure en deux étapes : LASSO pour sélectionner les bonnes composantes puis MCO pour estimer et ont observé que cette procédure de type Ante est significativement meilleure que le LASSO car elle le débiaise. Ces résultats sont cohérents avec ce que nous avons aussi constaté sur nos simulations, où les procédures Ante permettent d'améliorer significativement les prédictions faites à partir d'un estimateur Group LASSO. Ce résultat de Belloni *et al.* [16] inspire notre travail dans ce document. Nous nous focalisons sur le Group LASSO et les B-Splines. Pour ce travail, nous avons besoin de tenir plus précisément compte du fait que :

- Chaque composante additive f_j^* du modèle est approchée par une combinaison linéaire de fonctions de bases (les B-Splines), leur nombre pouvant croître avec le nombre d'observations n pour que l'approximation soit de bonne qualité. Nous avons donc besoin de

résultats nous assurant que cette approximation dans des bases de B-Splines n'est pas mauvaise. Stone [106] permet de justifier l'approximation.

- Nous appliquons le Group LASSO aux vecteurs des coefficients résultants de cette approximation puis nous ré-appliquons du Group LASSO ou des MCO ou des P-Splines pour améliorer la sélection (dans les cas des procédures de type Post) et l'estimation en débiaisant l'estimation par Group LASSO.

Nous avons constaté au cours des simulations que la première étape de nos procédures peut inclure des faux positifs ou des faux négatifs, quelque soit le critère de sélection de modèle utilisé (BIC, AIC ou GCV). Les simulations montrent qu'intégrer la seconde étape permet aux procédures de type Post d'améliorer les capacités de sélection de variables des procédures, et aux procédures de type Ante et Post d'améliorer fortement la capacité de prédiction des estimateurs par rapport à l'estimateur Group LASSO.

Dans la section 3.2, nous donnons les résultats qui prouvent la consistance en sélection et en estimation d'une de nos procédures. Les sections 3.4 et 3.5 contiennent les démonstrations des résultats annoncés. L'Annexe D contient des résultats de simulations.

3.2 Etude asymptotique des procédures

Notations

Par la suite, nous adoptons les notations suivantes :

$$\underline{B}_j(\cdot) = \{B_{j,k}^{q_j}(\cdot) | k = 1, \dots, K_j + q_j = m_j\},$$

avec $m_j = m_j(n) \rightarrow +\infty$ tel que $m_j = o(n)$. Avec ces deux hypothèses, nous nous assurons que m_j croît lorsque le nombre d'observations augmente (nécessaire pour avoir une bonne approximation des fonctions dans les bases de B-Splines) de manière lente. Nous poserons des hypothèses plus fortes sur m_j par la suite. Pour simplifier les notations, nous supposons que pour tout $j, j = 1, \dots, d$, les $m_j = m_n$ et les $\underline{B}_j(\cdot)$ sont identiques. Nous notons

$$\mathcal{B}_j = \{f : [0, 1] \rightarrow \mathbb{R} \mid f(\cdot) = \beta_0 + \sum_{k=1}^{m_n} \beta_k B_{j,k}^{q_j}(\cdot); \beta_k \in \mathbb{R}, \beta_0 \in \mathbb{R}\}.$$

Nous notons pour $1 \leq j \leq d$:

$$\tilde{\mathbf{x}}_{\mathbf{i}, \mathbf{B}_j} = \underline{B}_j^T(x_{ij}) - \frac{1}{n} \sum_{i=1}^n \underline{B}_j^T(x_{ij}) \in \mathbb{R}^{m_n},$$

et

$$\tilde{\mathbf{X}}_{\mathbf{i}} = [\tilde{\mathbf{x}}_{\mathbf{i}, \mathbf{B}_1}^T, \dots, \tilde{\mathbf{x}}_{\mathbf{i}, \mathbf{B}_d}^T]^T \in \mathbb{R}^{d \times m_n},$$

et pour $1 \leq j \leq d$,

$$\hat{\Sigma}_j = \frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{x}}_{\mathbf{i}, \mathbf{B}_j} \tilde{\mathbf{x}}_{\mathbf{i}, \mathbf{B}_j}^T,$$

de dimension $m_n \times m_n$. Nous notons la matrice de covariance empirique

$$\hat{\Sigma} = \frac{1}{n} \sum_{i=1}^n \tilde{\mathbf{X}}_{\mathbf{i}} \tilde{\mathbf{X}}_{\mathbf{i}}^T,$$

de dimension $m_n d \times m_n d$. Enfin, nous notons le vecteur des coefficients inconnus

$$\underline{\beta} = (\underline{\beta}_1^T, \dots, \underline{\beta}_d^T)^T \in \mathbb{R}^{dm_n},$$

et

$$\underline{\beta}_{B_j} = \underline{\beta}_j = (\beta_{j1}, \dots, \beta_{jm_n})^T.$$

La solution du Group LASSO standardisé consiste en :

$$\hat{\beta}_0 = \frac{1}{n} \sum_{i=1}^n Y_i \text{ et } \hat{f}_j(X_j) = \sum_{k=1}^{m_n} \hat{\beta}_{jk} (B_{j,k}^{q_j}(X_j) - \bar{B}_{j,k}^{q_j}),$$

où $\bar{B}_{j,k}^{q_j} = 1/n \sum_{i=1}^{m_n} B_{j,k}^{q_j}(X_{i,j})$ et

$$\hat{\underline{\beta}} = \arg \min_{\underline{\beta} \in \mathbb{R}^{dm_n}} \frac{1}{2} \sum_{i=1}^n (Y_i - \tilde{X}_i^T \underline{\beta})^2 + \sqrt{m_n} \lambda \sum_{j=1}^d \|\hat{\Sigma}^{1/2} \underline{\beta}_j\|_2, \quad (3.4)$$

où $\|\cdot\|_2$ désigne la norme euclidienne de $\mathbb{R}^{m_n}$.

Le Group LASSO est creux par groupe, ce qui signifie que si λ est suffisamment fort, il existe j telle que $\hat{\underline{\beta}}_j = \mathbf{0}$. Notons aussi que $\hat{\underline{\beta}}_j = \mathbf{0}$ si, et seulement si, $\hat{f}_j(\cdot) = 0$ (sauf si $\hat{\Sigma}_j$ est singulière). Nous avons utilisé une base commune pour toutes les covariables, ce qui n'est pas nécessaire pour utiliser le Group LASSO. Nos procédures Ante et Post dépendant de méthodes déjà existantes, l'implémentation est simple. Nous nous sommes concentrés sur le Group LASSO pour la première étape des procédures, nous aurions pu utiliser une autre pénalisation, qui tiendrait par exemple compte de la corrélation inter-groupe, s'il y en a.

Nous utilisons le Group LASSO standardisé qui a des bonnes propriétés asymptotiques quand le paramètre λ est choisi de manière optimale selon le critère de sélection adopté.

Nous notons

$$\|\mathbf{f}^*(\mathbf{x})\| = \|\mathbf{f}^*(\mathbf{x})\|_2 = \sqrt{\sum_{i=1}^n \left(\sum_{j=1}^d f_j^*(x_{i,j}) \right)^2},$$

et

$$\|\mathbf{f}^*\|_\infty = \max_{\mathbf{x}_i, i \in 1, \dots, n} |\mathbf{f}^*(\mathbf{x}_i)| = \max_{\mathbf{x}_i, i \in 1, \dots, n} \left| \sum_{j=1}^d f_j^*(x_{i,j}) \right|.$$

Nous noterons indifféremment $\|\cdot\|$ et $\|\cdot\|_2$ lorsque nous utilisons la norme euclidienne.

Par la suite, nous avons tout écrit avec $m_j = m = m(n)$ et la même base pour simplifier les notations mais les résultats obtenus par la suite sont généraux.

3.2.1 Hypothèses pour l'étude asymptotique

Meier *et al.* [89] donnent les conditions sous lesquelles une version modifiée de l'équation (3.3) permet de ne pas avoir de faux négatifs (voir le corollaire 1 de [89]). Pour que les conditions soient remplies, il faut que les composantes additives non nulles soient suffisamment éloignées de 0 au sens de la métrique que ces auteurs définissent. Cependant, ce résultat est insuffisant pour démontrer la consistance en sélection. De plus, dans les conditions du corollaire, il y a une condition de compatibilité, qui n'est pas vérifiable en pratique puisqu'elle nécessite la connaissance de S^* .

Nous énumérons et expliquons les différentes conditions qui sont nécessaires à la démonstration des résultats théoriques de nos estimateurs. Selon les résultats recherchés, il n'est pas nécessaire de supposer que toutes les conditions soient simultanément réalisées, mais il nous semble plus clair d'en énumérer la liste complète en premier.

Une référence importante dans le travail qui suit est l'article de Wang, Li et Leng [116]. Les résultats de Wang *et al.* [116] sont consacrés à l'analyse des modèles linéaires avec un nombre de covariables possibles tendant vers l'infini. Wang *et al.* [116] montrent que le BIC est un critère permettant d'identifier de manière consistante le vrai modèle lorsque le nombre de covariables potentielles tend vers l'infini pour les estimateurs, qu'ils soient pénalisés ou non. Nous nous sommes inspirés de leur travail pour l'appliquer au modèle additif décrit au début du chapitre.

Nous nous plaçons dans le même cadre que Huang *et al.* [67] au niveau des dimensions des plans d'expérience, c'est-à-dire :

- Le nombre de covariables candidates d peut éventuellement tendre vers l'infini
- Le nombre de composantes non nulles s^* ne dépend pas de n
- $s^* \leq D < n$ avec D fixé indépendant de n

Notons, pour $j = 1, \dots, d$,

$$\mathbf{U}_j = \begin{pmatrix} B_{j1}(x_{1j}) - \bar{B}_{j1} & \dots & B_{jm_n}(x_{1j}) - \bar{B}_{jm_n} \\ \vdots & & \vdots \\ B_{j1}(x_{nj}) - \bar{B}_{j1} & \dots & B_{jm_n}(x_{nj}) - \bar{B}_{jm_n} \end{pmatrix},$$

que nous pouvons ré-écrire :

$$\mathbf{U}_j = (\tilde{\mathbf{x}}_{1, \mathbf{B}_j} \dots \tilde{\mathbf{x}}_{n, \mathbf{B}_j})^T,$$

et posons $\mathbf{Y} = (Y_1, \dots, Y_n)^T$. Enfin, notons $\mathbf{U} = [\mathbf{U}_1 \dots \mathbf{U}_d]_{n \times dm_n}$.

Pour tout sous-modèle $S \subseteq \{1, \dots, d\}$, nous notons $\mathbf{U}_S$ la matrice extraite de $\mathbf{U}$ constituée des colonnes indexées par S , c'est-à-dire que

$$\mathbf{U}_S = [\mathbf{U}_j | j \in S].$$

De même, nous définissons $\boldsymbol{\beta}_S = (\beta_j)_{j \in S}$, $\hat{\boldsymbol{\beta}}_S = (\hat{\beta}_j)_{j \in S}$ (estimés sur S), et plus généralement $\boldsymbol{\Psi}_S = (\boldsymbol{\Psi}_j)_{j \in S}$.

Nous notons $\mathbf{1}_n$ le vecteur colonne de taille n dont toutes les composantes sont égales à 1. Pour simplifier les notations, nous ajoutons $\mathbf{1}_n / \sqrt{m_n}$ comme première colonne de $\mathbf{U}$ et de $\mathbf{U}_S$.

Posons

$$\underline{\mathbf{b}} = (\sqrt{m_n} \beta_0, \boldsymbol{\beta}^T)^T \text{ et } \underline{\mathbf{b}}_S = (\sqrt{m_n} \beta_0, \boldsymbol{\beta}_S^T)^T.$$

Nous pouvons alors ré-écrire matriciellement le critère MCO associé au sous-modèle S par l'équation (3.5) :

$$Q_S^{MCO}(\mathbf{b}_S) = \|\mathbf{Y} - \mathbf{U}_S \mathbf{b}_S\|^2. \quad (3.5)$$

Notons $\hat{\mathbf{b}}_S$ un vecteur minimisant (3.5), c'est-à-dire

$$\hat{\mathbf{b}}_S = \arg \min_{\mathbf{b}_S} Q_S^{MCO}(\mathbf{b}_S).$$

Notons que ce vecteur n'est pas unique si $n \leq \text{card}(S)m_n$.

Nous définissons (voir Wang *et al.* [116]) le critère BIC par l'équation (3.6) :

$$BIC(S) = \log(\|\mathbf{Y} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2) + \text{card}(S)m_n \frac{\log(nd)}{n}. \quad (3.6)$$

Pour être cohérent avec la définition usuelle du BIC, nous aurions dû utiliser $\log(\frac{\|\mathbf{Y} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2}{n})$ à la place de $\log(\|\mathbf{Y} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2)$, mais cela ne modifie pas le critère puisque nous cherchons S tel que $S = \arg \min BIC(S)$ et $\log \frac{\|\mathbf{Y} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2}{n} = \log(\|\mathbf{Y} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2) - \log(n)$.

Nous faisons les hypothèses suivantes :

1. **Hypothèse sur les vecteurs aléatoires** $(Y_i, \mathbf{X}_i)_{i=1, \dots, n} : (Y_i, \mathbf{X}_i)_{i=1, \dots, n}$ est un n -échantillon de vecteurs aléatoires $(Y_1, \mathbf{X}_1)$ de loi satisfaisant le modèle additif postulé dans le modèle (3.1).
2. **Hypothèse sur la loi des erreurs** : la loi conditionnelle de ϵ_1 sachant $\mathbf{X}_1$ est une loi $N(0, \sigma^2(\mathbf{X}_1))$ et $\sigma_1^2 \leq \sigma^2(\mathbf{X}_1) \leq \sigma^2$ p.s. avec σ^2 et σ_1^2 des constantes indépendantes de n (ce cas recouvre le cas classique où les erreurs sont indépendantes de $\mathbf{X}_1$).
3. **Hypothèse sur la loi de $\mathbf{X}_1$** : le support de $\mathbf{X}_1$ est compact. Pour alléger les notations, nous supposons que le support est contenu dans $[0, 1]^d$. De plus, les densités marginales de $\mathbf{X}_1$ sont bornées inférieurement et supérieurement par deux constantes strictement positives fixes. La densité de $\mathbf{X}_1$ est continue par rapport à la mesure de Lebesgue.
4. **Hypothèse sur la régularité des composantes additives** : pour tout $j \in S^*$, $f_j^* \in H_2^\nu([0, 1])$ avec ν entier supérieur ou égale à 2. L'inégalité de Morrey permet de dire que f_j^* est Holdérienne d'indice $> 1/2$.
5. **Hypothèses sur d, m_n , et s^*** :
 - (a) $m_n = An^{\frac{1}{2\nu+1}}$, avec A constante. Nous supposons donc que nous prenons des bases de splines de degré $q \geq 3$ (splines cubiques) pour approximer les composantes.
 - (b) $\frac{\log(d)}{n^{\frac{2\nu}{2\nu+1}}} \rightarrow 0$.
 - (c) $s^* \leq D \leq \min(d, \frac{n}{m_n})$.
 - (d) $\frac{m_n \log(nd)}{n \min_{j \in S^*} \|f_j^*\|^2} \rightarrow 0$, c'est-à-dire que $\frac{m_n \log(nd)}{n} = o(\min_{j \in S^*} \|f_j^*\|^2)$.

L'hypothèse 1 est une hypothèse classique en régression. L'hypothèse 2 est une hypothèse suffisante pour s'assurer une propriété de concentration de la loi normale (autour de sa moyenne et sachant $\mathbf{X}_1$) d'une variable aléatoire de la forme $\sup_{t \in T} \sum_{i=1}^n \epsilon_i t_i$, où T est un sous-ensemble de $\mathbb{R}^n$ borné, dénombrable et dépendant de $\mathbf{X}_1, \dots, \mathbf{X}_n$. L'hypothèse 3 est classique. L'hypothèse 4 est une hypothèse sur la régularité des composantes additives. L'hypothèse 5.a est imposée par le choix de Stone pour l'approximation des composantes additives dans des bases de B-Splines permettant de réaliser un compromis optimal entre le biais et la variance. L'hypothèse 5.b assure que le logarithme de la dimension d de covariables candidates puisse croître moins rapidement que $n^{\frac{2\nu}{2\nu+1}}$. Cette hypothèse nous permet donc d'affirmer que le nombre de covariables candidates croît mais n'explose pas. L'hypothèse 5.c permet de s'assurer de l'unicité de l'estimateur des MCO et est nécessaire pour des raisons d'identifiabilité. Enfin, l'hypothèse 5.d est suffisante pour nous permettre de démontrer des résultats de consistance en sélection et en estimation. Cette hypothèse est intéressante, parce qu'elle n'oblige pas la norme euclidienne de chaque composante à être strictement supérieure à une constante strictement positive indépendante de n .

3.2.2 Consistance asymptotique en sélection du BIC appliqué à l'estimateur des MCO

Dans cette sous-section, nous montrons que sous les conditions 1 à 5, le critère BIC défini par l'équation (3.6) est consistant en sélection de variables, c'est-à-dire que

$$P(\hat{S} = S^*) \xrightarrow{n \rightarrow \infty} 1.$$

Nous considérons deux cas. Dans le premier cas, nous supposons que nous ajustons un modèle S tel que $S^* \not\subseteq S$. Dans le second cas, nous supposons que nous ajustons un modèle S tel que $S^* \subseteq S$.

3.2.2.1 Cas où $S^* \not\subseteq S$

Considérons le cas où $S^* \not\subseteq S$. Notons $\hat{\mathbf{b}}_S$ et $\hat{\mathbf{b}}_{S^*}$ les estimateurs des MCO sous les modèles S et S^* , c'est-à-dire

$$\hat{\mathbf{b}}_S = \arg \min_{\mathbf{b}} \|\mathbf{Y} - \mathbf{U}_S \mathbf{b}\|^2 \text{ et } \hat{\mathbf{b}}_{S^*} = \arg \min_{\mathbf{b}} \|\mathbf{Y} - \mathbf{U}_{S^*} \mathbf{b}\|^2.$$

Nous avons supposé que $\text{card}(S) \leq D$. L'hypothèse 5.c nous assure que $m_n D < n$.
Notons alors

$$\bar{S} = S^* \cup S,$$

et notons aussi $\tilde{\mathbf{b}}_S$ le vecteur de dimension $\text{card}(\bar{S})m_n + 1$ coïncidant avec $\hat{\mathbf{b}}_S$ pour les coefficients appartenant à S et nul pour les autres composantes, c'est-à-dire

$$\tilde{\mathbf{b}}_S = \left(\bigcup_{j \in \bar{S}} \tilde{\mathbf{b}}_{S,j}^T \right) \text{ avec } \tilde{\mathbf{b}}_{S,j} = \hat{\mathbf{b}}_{S,j} \text{ si } j \in S, \mathbf{0} \text{ sinon.}$$

De même, nous notons

$$\tilde{\mathbf{b}}_{S^*} = \left(\bigcup_{j \in \bar{S}} \tilde{\mathbf{b}}_{S^*,j}^T \right) \text{ avec } \tilde{\mathbf{b}}_{S^*,j} = \hat{\mathbf{b}}_{S^*,j} \text{ si } j \in S^*, \mathbf{0} \text{ sinon,}$$

et

$$\mathbf{b}_{\bar{S}}^* = \left(\bigcup_{j \in \bar{S}} \mathbf{b}_{\bar{S},j}^{*,T} \right) \text{ avec } \mathbf{b}_{\bar{S},j}^* = \mathbf{b}_{S^*,j}^* \text{ si } j \in S^*, \mathbf{0} \text{ sinon.}$$

Avec ces notations, nous prolongeons par des 0 lorsque cela est nécessaire les divers vecteurs considérés pour que les produits matriciels aient un sens.

Nous remarquons que

$$\|\mathbf{Y} - \mathbf{U}_{\bar{S}} \tilde{\mathbf{b}}_S\|^2 = \|\mathbf{Y} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2,$$

et de plus,

$$BIC(S) - BIC(S^*) = \log \left(1 + \frac{\|\mathbf{Y} - \mathbf{U}_{\bar{S}} \tilde{\mathbf{b}}_S\|^2/n - \|\mathbf{Y} - \mathbf{U}_{\bar{S}} \tilde{\mathbf{b}}_{S^*}\|^2/n}{\|\mathbf{Y} - \mathbf{U}_{\bar{S}} \tilde{\mathbf{b}}_{S^*}\|^2/n} \right) + (\text{card}(S) - \text{card}(S^*)) m_n \frac{\log(nd)}{n}. \quad (3.7)$$

Nous montrons que $BIC(S) - BIC(S^*) > 0$ avec probabilité tendant vers 1. Plus précisément,

Théorème 3.1 Soit S tel que $S^* \not\subseteq S$. Alors,

$$P\left(\min_{S^* \not\subseteq S, \text{card}(S) \leq M} BIC(S) - BIC(S^*) > 0 \right) \xrightarrow{n \rightarrow \infty} 1,$$

où M constante.

En particulier, une des conséquences est qu'il n'y a pas de faux négatifs avec probabilité tendant vers 1 lorsque le vrai plan d'expérience est candidat.

Pour démontrer le Théorème 3.1, nous étudions d'abord le terme donné dans l'équation (3.8) :

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\tilde{\mathbf{S}}}\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*}\|^2. \quad (3.8)$$

Nous pouvons montrer grâce à des calculs classiques de développement de la norme euclidienne et en utilisant le fait que $\underline{\mathbf{Y}} = \underline{\mathbf{f}}^*(\underline{\mathbf{X}}) + \underline{\epsilon}$ l'identité (3.9) suivante :

$$\begin{aligned} \|\underline{\mathbf{Y}} - \mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\tilde{\mathbf{S}}}\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*}\|^2 \\ = -2\underline{\epsilon}^T \mathbf{U}_{\tilde{\mathbf{S}}} (\tilde{\mathbf{b}}_{\tilde{\mathbf{S}}} - \tilde{\mathbf{b}}_{\mathbf{S}^*}) + 2 \left(\mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*} - \underline{\mathbf{f}}^*(\underline{\mathbf{X}}) \right)^T \mathbf{U}_{\tilde{\mathbf{S}}} (\tilde{\mathbf{b}}_{\tilde{\mathbf{S}}} - \tilde{\mathbf{b}}_{\mathbf{S}^*}) \\ + \|\mathbf{U}_{\tilde{\mathbf{S}}} (\tilde{\mathbf{b}}_{\tilde{\mathbf{S}}} - \tilde{\mathbf{b}}_{\mathbf{S}^*})\|^2, \end{aligned} \quad (3.9)$$

où $[\underline{\mathbf{f}}^*(\underline{\mathbf{X}})]_{ij} = \beta_0^* + \sum_j f_j^*(X_{i,j})$.

Le premier terme du membre de droite de l'équation (3.9) correspond à la multiplication de l'erreur propre au modèle par la différence des prédictions faites lorsque nous utilisons les deux plans d'expérience différents. Le second terme de droite correspond au coût de l'approximation dans des bases de B-Splines tronquées et de l'utilisation de MCO pour estimer la vraie fonction de régression. On multiplie ce coût par la différence des prédictions selon le plan d'expériences utilisé. Le dernier terme de droite représente la norme L2 au carré de cette différence.

Stone [106] a montré que lorsque $\mathbf{S}^*$ est connu, nous avons :

$$\|\hat{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{b}_{\mathbf{S}^*}^*\| = O\left(\frac{m_n}{\sqrt{n}} + m_n^{-\nu+1/2}\right),$$

et de plus, lorsque $f_j^* \in H_2^\nu$,

$$\|f_j^* - \sum_{k=1}^{m_n} b_{jk}^* \tilde{B}_{j,k}\|_\infty = O(m_n^{-\nu}).$$

La première égalité permet de contrôler l'erreur commise parce que nous utilisons des estimateurs MCO pour estimer les coefficients associés aux bases de B-Splines, alors que la seconde permet de contrôler uniformément le biais pour approcher les composantes additives par des projections dans des bases de B-Splines.

Le contrôle de la différence donnée dans le terme du membre de gauche de l'équation (3.9) est obtenu en étudiant séparément les trois termes du membre de droite de l'équation (3.9). Nous commençons par l'analyse de troisième terme. Le Lemme 3.1 en donne une minoration.

Lemme 3.1

$$\|\mathbf{U}_{\tilde{\mathbf{S}}} (\tilde{\mathbf{b}}_{\tilde{\mathbf{S}}} - \tilde{\mathbf{b}}_{\mathbf{S}^*})\|^2 \geq A_2 \frac{n}{m_n} v_n^2, \quad (3.10)$$

où A_2 une constante et

$$v_n = A_3 \left(\sqrt{m_n} \min_{j \in \mathbf{S}^*} \|f_j^*\| - \frac{m_n}{\sqrt{n}} - m_n^{-\nu+1/2} \right), \quad (3.11)$$

avec A_3 une constante.

Nous constatons que nous pouvons minorer le troisième terme par une suite indépendante du plan d'expérience S . De plus, nous remarquons que

$$\frac{m_n}{\sqrt{n}} + m_n^{-\nu+1/2} = 2An^{\frac{-\nu+1/2}{2\nu+1}} \rightarrow 0.$$

Nous montrons également dans la section 3.4 le Lemme 3.2, qui établit que le second terme du membre de droite de l'équation (3.9) est négligeable par rapport au troisième terme.

Lemme 3.2

$$|2(\mathbf{U}_{\tilde{S}}\tilde{\mathbf{b}}_{S^*} - \mathbf{f}^*(\mathbf{X}))^T \mathbf{U}_{\tilde{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})| \leq C_4 \sqrt{nm_n^{-2\nu}} \sqrt{\frac{n}{m_n}} |v_n|,$$

avec $C_4 > 0$ constante. L'ordre de ce terme est négligeable par rapport au troisième terme de l'équation (3.9).

Il reste à étudier le premier terme du membre de droite de l'équation (3.9), que nous ne pouvons pas minorer par un terme strictement positif. Cependant, nous montrons que nous pouvons minorer la somme du premier terme et du troisième terme de droite par une suite strictement positive. Pour prouver ce résultat, nous établissons des résultats à propos du coût d'approximation des composantes du modèle additif dans des bases de B-Splines (Proposition 3.1 de la section 3.4), ainsi que sur l'expression $\underline{\epsilon}^T \Pi_{\tilde{S}} \underline{\epsilon}$ (Proposition 3.2 de la section 3.4). Nous établissons aussi la Proposition 3.3 de la section 3.4, qui est centrale pour la suite de la démonstration et montre que $\|\underline{\mathbf{Y}} - \mathbf{U}_{\tilde{S}}\hat{\mathbf{b}}_{\tilde{S}}\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{\tilde{S}}\mathbf{b}_{\tilde{S}}^*\|^2 = O_p(nm_n^{-2\nu}) + o_p(m_n \log(nd))$ lorsque $S^* \subseteq \tilde{S}$.

Puisque $S^* \subseteq \tilde{S}$, nous pouvons écrire

$$\begin{aligned} \underline{\mathbf{Y}} &= \mathbf{f}^*(\mathbf{X}) + \underline{\epsilon} \\ &= \mathbf{f}^*(\mathbf{X}) - \mathbf{U}_{\tilde{S}}\mathbf{b}_{\tilde{S}}^* + \mathbf{U}_{\tilde{S}}\mathbf{b}_{\tilde{S}}^* + \underline{\epsilon} \\ &= B + \underline{\epsilon} + \mathbf{U}_{\tilde{S}}\mathbf{b}_{\tilde{S}}^*, \end{aligned}$$

avec $B = \mathbf{f}^*(\mathbf{X}) - \mathbf{U}_{\tilde{S}}\mathbf{b}_{\tilde{S}}^*$.

De plus, nous notons $\Pi_{\tilde{S}}$ le projecteur orthogonal sur l'espace engendré par les colonnes de $\mathbf{U}_{\tilde{S}}^T$:

$$\Pi_{\tilde{S}} = \mathbf{U}_{\tilde{S}}(\mathbf{U}_{\tilde{S}}^T \mathbf{U}_{\tilde{S}})^{-1} \mathbf{U}_{\tilde{S}}^T.$$

Le lemme suivant permet d'affirmer que la somme du premier et du troisième terme du membre de droite de l'équation (3.9) est minorée par une suite strictement positive.

Lemme 3.3 *Sous nos hypothèses,*

$$-2\underline{\epsilon}^T \mathbf{U}_{\tilde{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*}) \geq -4\underline{\epsilon}^T \Pi_{\tilde{S}} \Pi_{\tilde{S}} \underline{\epsilon} - \frac{1}{4} \|\mathbf{U}_{\tilde{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})\|^2. \quad (3.12)$$

Nous avons donc

$$-2\underline{\epsilon}^T \mathbf{U}_{\tilde{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*}) + \|\mathbf{U}_{\tilde{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})\|^2 \geq \frac{3}{4} \|\mathbf{U}_{\tilde{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})\|^2 + o_p(m_n \log(nd)) \geq R_1 \frac{n}{m_n} v_n^2,$$

avec R_1 constante.

Grâce aux trois lemmes précédents, nous obtenons que

$$\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\tilde{\mathbf{S}}}\|^2 - \|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\tilde{\mathbf{S}}^*}\|^2 \geq N_1 \frac{n}{m_n} v_n^2 + o_p\left(\frac{n}{m_n} v_n^2\right), \quad (3.13)$$

avec N_1 constante.

Les hypothèses sur les normes des composantes additives impliquent en particulier que :

$$\frac{v_n^2}{m_n} \geq \min \|f_j^*\|^2 + o(\min \|f_j^*\|^2).$$

Ceci permet d'établir le Lemme 3.4 :

Lemme 3.4 *Les suites*

$$\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}^*} \mathbf{b}_{\mathbf{S}^*}\|^2/n,$$

et

$$\inf_{S^* \subseteq S \mid \text{card}(S) \leq D} \|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}} \hat{\mathbf{b}}_{\mathbf{S}}\|^2/n,$$

sont des suites bornées inférieurement par une constante strictement positive avec une probabilité tendant vers 1.

D'après ce lemme, $\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*}\|/n$ est bornée inférieurement. L'équation (3.13) et les résultats précédents montrent qu'il existe une constante N_2 telle que, avec une probabilité tendant vers 1,

$$\frac{\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\tilde{\mathbf{S}}}\|^2/n - \|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\tilde{\mathbf{S}}^*}\|^2/n}{\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\tilde{\mathbf{S}}^*}\|^2/n} \geq N_2 \frac{v_n^2}{m_n} > \min \|f_j^*\|^2 + o(\min \|f_j^*\|^2) > 0.$$

Comme il est montré dans la section 3.4, nous pouvons conclure la démonstration du Théorème 3.1 en utilisant $\log(x) \leq \min(x/2, \log(2))$.

3.2.2.2 Cas où $S^* \subseteq S$

Nous montrons ensuite le résultat du Théorème 3.2, qui traite des faux positifs.

Théorème 3.2

$$P\left(\min_{S^* \subseteq S \mid \text{card}(S) \leq M} BIC(S^*) - BIC(S) < 0\right) \xrightarrow{n \rightarrow \infty} 1.$$

Notons, de la même manière que dans le paragraphe précédent, par $\underline{\mathbf{b}}_{\mathbf{S}}^*$ le vecteur $\underline{\mathbf{b}}_{\mathbf{S}^*}^*$ prolongé par des 0 pour les j tels que $j \in S$ et $j \notin S^*$.

Nous étudions l'expression :

$$BIC(S^*) - BIC(S) = \log \left(\frac{\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*}\|^2}{\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}} \hat{\mathbf{b}}_{\mathbf{S}}\|^2} \right) + (\text{card}(S^*) - \text{card}(S)) m_n \frac{\log(nd)}{n}.$$

Par définition de l'estimateur MCO,

$$\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*}\|^2 \leq \|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}^*} \underline{\mathbf{b}}_{\mathbf{S}^*}^*\|^2.$$

Donc, nous avons

$$BIC(S^*) - BIC(S) \leq \log \left(1 + \frac{\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}^*} \underline{\mathbf{b}}_{\mathbf{S}^*}^*\|^2 - \|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}} \hat{\mathbf{b}}_{\mathbf{S}}\|^2}{\|\underline{\mathbf{Y}} - \underline{\mathbf{U}}_{\mathbf{S}} \hat{\mathbf{b}}_{\mathbf{S}}\|^2} \right) + (\text{card}(S^*) - \text{card}(S)) m_n \frac{\log(nd)}{n}.$$

Nous montrons dans la section 3.4 le lemme suivant :

Lemme 3.5

$$\frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \underline{\mathbf{b}}_{\mathbf{S}^*}^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \hat{\underline{\mathbf{b}}}_{\mathbf{S}}\|^2}{\|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \hat{\underline{\mathbf{b}}}_{\mathbf{S}}\|^2} \in [0, 1], \quad (3.14)$$

avec une probabilité tendant vers 1.

Dans la démonstration du Lemme 3.5, nous montrons que

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \underline{\mathbf{b}}_{\mathbf{S}^*}^*\|^2/n - \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \hat{\underline{\mathbf{b}}}_{\mathbf{S}}\|^2/n = \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \underline{\mathbf{b}}_{\mathbf{S}}^*\|^2/n - \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \hat{\underline{\mathbf{b}}}_{\mathbf{S}}\|^2/n \xrightarrow{P} 0,$$

ce qui montre que les erreurs quadratiques moyennes obtenues lorsque les “vrais” coefficients sont utilisés tendent à prendre les mêmes valeurs que les erreurs quadratiques moyennes des estimateurs MCO (qui sont minimales).

Dans la section 3.4, nous utilisons le Lemme 3.5 et le développement limité en 0 du logarithme pour démontrer le Théorème 3.2.

Le Théorème 3.1 nous permet d’affirmer que si $\mathbf{S}^*$ est candidat, tous les ensembles $\mathbf{S}$ de cardinal fini ne contenant pas $\mathbf{S}^*$ ont un BIC supérieur avec une probabilité tendant vers 1. Ce théorème permet ainsi de se prémunir des faux négatifs. Ce résultat et le Théorème 3.2 permettent d’affirmer qu’il n’y a pas de faux positifs avec une probabilité tendant vers 1.

3.2.2.3 Théorème de consistance du BIC

Nous pouvons alors énoncer le Théorème 3.3 qui est analogue au Théorème 3 de Wang *et al.* [116] et qui permet d’affirmer la consistance du BIC.

Théorème 3.3 *Sous les conditions 1 à 5, le critère BIC défini par (3.6) est consistant en sélection de variables, c’est-à-dire*

$$P(\hat{\mathbf{S}} = \mathbf{S}^*) \xrightarrow{n \rightarrow \infty} 1.$$

Le BIC défini par l’équation (3.6) est donc consistant en termes de sélection de variables. Ceci n’est pas suffisant pour prouver la consistance de notre procédure d’estimation. En effet, nous venons de démontrer que dans le cas où nous estimons à l’aide de MCO un modèle additif dont les composantes sont approchées dans une base de B-Splines, si le vrai modèle est parmi les modèles candidats, alors il est sélectionné avec probabilité tendant vers 1 lorsque le BIC est utilisé comme critère de sélection.

Toutefois, comme nous venons de le voir, le critère BIC que nous avons étudié n’est pas le critère BIC de l’estimateur que nous avons défini dans l’équation (3.4), car il ne tient pas compte de la pénalisation.

3.2.3 Etude de la procédure Post1**3.2.3.1 Consistance en sélection de la procédure Post1**

Le Théorème 3.3 permet d’affirmer la consistance en sélection du BIC lorsqu’un estimateur MCO est utilisé. Nous montrons ensuite que le BIC appliqué à l’estimateur défini en (3.4) est consistant.

Sous nos hypothèses, Kato [71] établit la Proposition 3.4 (voir la section 3.5), qui permet de simplifier l’équation (3.4) et d’affirmer que nous pouvons remplacer $\hat{\Sigma}_j$ par la matrice identité

dans l'expression (3.4). De plus, le centrage des B-Splines n'a pas d'impact sur les résultats asymptotiques. Nous pouvons ainsi ré-écrire l'équation (3.4) sous la forme de l'équation (3.15) :

$$\tilde{\mathbf{b}}_\lambda = \arg \min_{\mathbf{b} \in \mathbf{R}^{dm_n+1}} \|\mathbf{Y} - \mathbf{U}^T \mathbf{b}\|^2 + \sqrt{m_n} \lambda \sum_{j=1}^d \|\mathbf{b}_j\|_2. \quad (3.15)$$

Il se pose la question du choix du paramètre λ . Le paramètre de lissage λ est choisi sur une grille de valeurs positives en sélectionnant celui qui minimise le BIC pénalisé.

Comme Huang *et al.* [67] l'ont constaté et utilisé pour le LASSO adaptatif, ce choix de λ permet de réduire le nombre de composantes additives et d'estimer un nombre $\text{card}(\hat{S})$ de composantes avec une très faible probabilité de faux négatifs.

Par la suite, nous précisons quelques notations. Nous notons

$$S_\lambda = \{j; \|\tilde{\mathbf{b}}_{\lambda,j}\| \neq 0\}.$$

De même, nous notons

$$BIC(\lambda) = \log(\|\mathbf{Y} - \mathbf{U} \tilde{\mathbf{b}}_\lambda\|^2) + \text{card}(S_\lambda) m_n \frac{\log(nd)}{n}.$$

De plus, nous notons

$$\hat{\mathbf{b}}_\lambda = (\tilde{\mathbf{b}}_{\lambda,j})_{j \in S_\lambda}.$$

À λ fixé, nous justifions dans la Proposition 3.5 (voir la section 3.5) qu'il est équivalent de résoudre le problème d'optimisation en considérant $\mathbf{U}$ ou $\mathbf{U}_{S_\lambda}$ (à la dimension des vecteurs près). Ce résultat est important, puisqu'il permet de travailler sur le plan d'expérience $\mathbf{U}_{S_\lambda}$, qui est de dimension finie par hypothèse, et non sur $\mathbf{U}$.

Nous avons donc

$$\begin{aligned} BIC(\lambda) &= \log(\|\mathbf{Y} - \mathbf{U}_{S_\lambda} \hat{\mathbf{b}}_\lambda\|^2) + \text{card}(S_\lambda) m_n \frac{\log(nd)}{n} \\ &= \log(\|\mathbf{Y} - \mathbf{U} \tilde{\mathbf{b}}_\lambda\|^2) + \text{card}(S_\lambda) m_n \frac{\log(nd)}{n}. \end{aligned}$$

Nous énonçons le théorème affirmant la consistance du BIC pénalisé.

Théorème 3.4 *Sous nos hypothèses,*

$$\text{soit } \lambda_n \sim \sqrt{n \log(m_n^2 d) / m_n},$$

alors

$$P(S_{\lambda_n} = S^*) \rightarrow 1,$$

et pour tout λ tel que $S_\lambda \neq S^$ et $\text{card}(S_\lambda) \leq M$,*

$$P(BIC(\lambda) - BIC(\lambda_n) > 0) \rightarrow 1.$$

Le Théorème 3.4 permet d'affirmer la consistance en sélection de l'estimateur issu de la procédure Post1. Pour démontrer ce théorème, nous montrons qu'à un terme négligeable près, il est équivalent de minimiser le BIC pénalisé et le BIC étudié dans la section précédente. La première étape est de donner une solution au programme d'optimisation associé au Group LASSO.

Soit λ_n tel que $S_{\lambda_n} = S^*$. Alors nous pouvons facilement montrer que

$$\hat{\mathbf{b}}_{\lambda_n} = (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} (\mathbf{U}_{S^*}^T \mathbf{Y} + \mathbf{r}),$$

avec

$$\begin{aligned}\underline{\mathbf{r}} &= -\lambda_n \sqrt{m_n} \frac{\partial}{\partial \underline{\mathbf{b}}} \sum_{j \in S^*} \|\underline{\mathbf{b}}_j\|_{\underline{\mathbf{b}}=\hat{\mathbf{b}}_{\lambda_n}} \\ &= -\lambda_n \sqrt{m_n} \left(0 \dots \left(\frac{\hat{\mathbf{b}}_{\lambda_n, \mathbf{j}}^T}{\|\hat{\mathbf{b}}_{\lambda_n, \mathbf{j}}\|} \right)_{j \in S^*} \right)^T.\end{aligned}$$

Le terme $\underline{\mathbf{r}}$ est inconnu, puisqu'il dépend des dérivées partielles des coefficients. Nous pouvons facilement montrer que

$$\|\underline{\mathbf{r}}\|^2 = O(\lambda_n^2 m_n).$$

$\underline{\mathbf{r}}$ est le terme qui différencie l'estimateur Group LASSO de l'estimateur MCO sur le plan d'expérience $\mathbf{U}_{S^*}$ et de coefficient de pénalisation λ_n . En effet,

$$\hat{\mathbf{b}}_{\lambda_n} = \hat{\mathbf{b}}_{S^*} + (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}.$$

Nous énonçons ensuite les résultats qui montrent que si $\lambda_n \sim \sqrt{n \log(m_n^2 d)/m_n}$, alors

$$P(S_{\lambda_n} = S^*) \rightarrow 1.$$

La Proposition 3.19 (voir la section 3.5) permet d'affirmer que $S_{\lambda_n} = S^*$ si

$$\begin{cases} \|\underline{\mathbf{b}}_j^*\| - \|\tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\| < \|\underline{\mathbf{b}}_j^*\| \text{ si } j \in S^*, \\ \|\mathbf{U}_j^T (\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n, S^*})\| \leq \lambda_n \sqrt{m_n} \text{ si } j \notin S_{\lambda_n}. \end{cases}$$

Nous voulons montrer qu'avec $\lambda_n \sim \sqrt{n \log(m_n^2 d)/m_n}$, $P(S_{\lambda_n} = S^*) \rightarrow 1$. Nous montrons plutôt que $P(S_{\lambda_n} \neq S^*) \rightarrow 0$. D'après la Proposition 3.6, il est facile de montrer que

$$P(S_{\lambda_n} \neq S^*) \leq P(\text{il existe } j \in S^*, \|\underline{\mathbf{b}}_j^* - \tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\| \geq \|\underline{\mathbf{b}}_j^*\|) + P(\text{il existe } j \notin S^*, \|\mathbf{U}_j^T (\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n, S^*})\| > \lambda_n \sqrt{m_n}).$$

Par la suite, nous montrons que si $\lambda_n \sim \sqrt{n \log(m_n^2 d)/m_n}$, alors

$$P(\text{il existe } j \in S^*, \|\underline{\mathbf{b}}_j^* - \tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\| \geq \|\underline{\mathbf{b}}_j^*\|) \rightarrow 0 \text{ et } P(\text{il existe } j \notin S^*, \|\mathbf{U}_j^T (\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n, S^*})\| > \lambda_n \sqrt{m_n}) \rightarrow 0.$$

Lemme 3.6

$$P(\text{il existe } j \in S^*, \|\underline{\mathbf{b}}_j^* - \tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\| \geq \|\underline{\mathbf{b}}_j^*\|) \rightarrow 0.$$

Ce lemme, démontré dans la section 3.5, permet d'affirmer l'absence de faux négatifs avec une probabilité tendant vers 1. Les résultats de Stone sont nécessaires pour montrer ce résultat.

Lorsque $\lambda_n \sim \sqrt{n \log(m_n^2 d)/m_n}$, il n'y a pas de faux négatifs avec une probabilité tendant vers 1 pour le Group LASSO. Nous travaillons ensuite sur l'absence de faux positifs. Pour montrer ce résultat, nous supposons que pour tout $j \neq j'$, $\|\mathbf{U}_j^T \mathbf{U}_{j'}\| = o_P(\frac{n}{m_n})$. Les hypothèses de ce type sont classiques en sélection de variables. Elles signifient que la matrice de plan d'expérience satisfait une hypothèse d'isométrie restreinte. Cette propriété requiert essentiellement que chaque couple $\mathbf{U}_j^T \mathbf{U}_{j'}$ soit semblable à un système orthonormal. Nous pouvons facilement trouver des exemples où cette hypothèse est vérifiée (par exemple, dans le cas de covariables indépendantes et suivant une loi uniforme).

Le Lemme 3.7 montre qu'il ne peut exister de faux positifs sous nos hypothèses avec une probabilité tendant vers 1 lorsque $\lambda_n \sim \sqrt{n \log(m_n^2 d)/m_n}$.

Lemme 3.7

$$P(\text{il existe } j \notin S^*, \|\mathbf{U}_j^T(\mathbf{Y} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n, S^*})\| > \lambda_n \sqrt{m_n}) \rightarrow 0.$$

La première étape de la démonstration consiste à se ramener au cas des EMCO, ce qui introduit une première source d'erreur, qui est contrôlée par un terme négligeable par rapport à $\lambda_n \sqrt{m_n}$ grâce à l'hypothèse qui suppose que pour tout $j \neq j'$, $\|\mathbf{U}_j^T \mathbf{U}_{j'}\| = o_P(\frac{n}{m_n})$. Ensuite, nous cherchons à majorer par trois suites négligeables par rapport à $\lambda_n \sqrt{m_n}$ les trois sources d'erreurs suivantes :

- celles dues aux résidus du modèle ;
- celles dues à l'approximation du modèle dans des bases de B-Splines ;
- celles dues aux erreurs commises par l'EMCO pour estimer les coefficients approchant le modèle.

Nous avons exhibé un λ_n tel que, sous nos hypothèses, le Group LASSO soit consistant en sélection.

Dans la section 3.5, nous montrons la Proposition 3.7, qui permet de montrer $\|\mathbf{U}_{S^*}(\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r}\|$ est bornée par $O(\sqrt{m_n \log(m_n^2 d)})$. La démonstration se fait principalement grâce au Lemme 3 de [67]. La Proposition 3.8 montrée dans la section 3.5 permet de donner les expressions de $\tilde{\mathbf{b}}_{\lambda_n} - \mathbf{b}_{S^*}^*$ et $\mathbf{Y} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n}$ en fonction de $\mathbf{r}$, de $\underline{\epsilon}$, du plan d'expérience et du biais d'approximation dans des bases de B-Splines.

La suite de la démonstration consiste à montrer que son BIC pénalisé est minimal avec une probabilité tendant vers 1. Nous montrons que la différence des erreurs quadratiques de l'estimateur associé au Group LASSO et de l'estimateur par MCO lorsque le vrai plan d'expérience a été sélectionné peut être majorée. Nous montrons d'abord, grâce à des développements classiques de la norme euclidienne, que cette différence est dépendante de $\mathbf{r}$ lorsque le vrai plan d'expérience est sélectionné.

$$\|\mathbf{Y} - \mathbf{U} \tilde{\mathbf{b}}_{\lambda_n}\|^2 - \|\mathbf{Y} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2 = \|\mathbf{U}_{S^*}(\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r}\|^2 - 2(\mathbf{Y} - \Pi_{S^*} \mathbf{Y})^T (\mathbf{U}_{S^*}(\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r}).$$

Le Lemme 3.8 permet de montrer que cette différence ne croît pas trop vite.

Lemme 3.8

$$\|\mathbf{Y} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2 - \|\mathbf{Y} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2 = o_P(m_n \sqrt{\log(m_n^2 d) \log(nd)}).$$

Maintenant que nous avons pu majorer la différence entre l'erreur de l'estimateur Group LASSO et de l'EMCO lorsque le vrai plan d'expérience est sélectionné, nous pouvons majorer la différence entre les deux BIC associés à ces estimateurs.

Nous avons $\text{card}(S_{\lambda_n}) = \text{card}(S^*)$ et donc

$$BIC(\lambda_n) - BIC(S^*) = \log(\|\mathbf{Y} - \mathbf{U} \tilde{\mathbf{b}}_{\lambda_n}\|^2) - \log(\|\mathbf{Y} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2).$$

Le Lemme 3.9 permet de majorer $BIC(\lambda_n) - BIC(S^*)$.

Lemme 3.9

$$BIC(\lambda_n) - BIC(S^*) = o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n}) + o(o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n})),$$

et $m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n}$ négligeable par rapport à $\min(\log(2), v_n^2/m_n)$.

Nous suivons le même genre de raisonnement que dans la sous-section précédente pour démontrer le Lemme 3.9.

Nous avons maintenant tous les résultats nécessaires pour démontrer le Théorème 3.4. Pour cela, nous utilisons un raisonnement similaire à Wang *et al.* [116] pour montrer que si $\lambda_n \sim \sqrt{n \log(m_n^2 d)/m_n}$, alors

$$P(BIC(\lambda) - BIC(\lambda_n) > 0) \rightarrow 1.$$

Nous venons ainsi d'énoncer les résultats permettant de montrer que le modèle choisi par le BIC pénalisé est consistant en sélection de variables. Il reste à démontrer l'optimalité en estimation de la procédure.

3.2.3.2 Consistance en estimation

Nous rappelons le Théorème 4.2 de Kato [72] utilisé pour démontrer l'optimalité de la procédure Post1Bic.

Théorème 3.5 *Sous les hypothèses 1 à 4 et 5.b, et avec des restrictions sur les bases utilisées à la première étape, si de plus $(s^*)^2 m_n \log(\max(n, d))/n \rightarrow 0$, si $\max_{card(S) \leq s, \alpha \in \mathbb{S}_T^{dm_n-1}} \|\hat{\Sigma}^{1/2} \mathbf{b}\| \leq O(1)$ tel que $s/s^* \rightarrow +\infty$, avec $\mathbb{S}^{dm_n-1}$ la sphère unité sur $\mathbb{R}^{dm_n}$ et $\mathbb{S}_T^{dm_n-1} = \mathbb{S}^{dm_n-1} \cap \{\alpha \in \mathbb{R}^{dm_n} | \alpha_{T^c} = \mathbf{0}\}$. Et si $\max_{\alpha \in \mathbb{S}^{dm_n-1} \cap \mathbb{C}} \|\hat{\Sigma}^{1/2} \mathbf{b}\| \geq O(1)$, avec $\mathbb{C} = \{\alpha \in \mathbb{R}^{dm_n} | \sum_{j \in (S^*, c)} \|\alpha_j\| \leq 21 \sum_{j \in (S^*)} \|\alpha_j\|\}$. Si de plus $\|g^*\| \leq O(1)s^*$. Alors il existe une constante A_f telle que*

$$1/n \|\mathbf{f}^*(\mathbf{X}) - \hat{\mathbf{f}}(\mathbf{X})\|^2 \leq_P \frac{s^* m_n}{n^2} A_f n (1 + \sqrt{\log(d/m_n)})^2,$$

où $\leq_P$ signifie inférieur ou égale avec une probabilité 1.

Les bases de B-Splines vérifient bien les conditions posées par Kato [72]. Or

$$\begin{aligned} \frac{s^* m_n}{n^2} A_f n (1 + \sqrt{\log(d/m_n)})^2 &= A_s \frac{m_n}{n} (1 + \sqrt{\log(d/m_n)})^2 \text{ (avec } A_s \text{ constante)} \\ &= A_{s_2} \frac{(1 + \sqrt{\log(d/m_n)})^2}{n^{2\nu/(2\nu+1)}} \text{ (avec } A_{s_2} \text{ constante)} \\ &\rightarrow 0 \text{ d'après l'hypothèse 5.b.} \end{aligned}$$

L'optimalité de la procédure Post1Bic découle donc du Théorème 3.5 de Kato [72] qui montre que le modèle B-Splines pénalisé avec une pénalisation de l'ordre $\sqrt{n \log(m_n^2 d)/m_n}$ conduit à une estimation consistante de la fonction de régression avec une erreur empirique quadratique moyenne de l'ordre

$$O(1) \frac{s^* m_n}{n} (1 + \sqrt{\log(d/m_n)})^2 \rightarrow 0.$$

Cela justifie entièrement la procédure Post1Bic.

3.3 Conclusion

En guise de conclusion, nous donnons ici une synthèse de la contribution.

L'intérêt des estimateurs en plusieurs étapes est illustré par à une simulation de Meier [89] qui montre que l'estimateur Group LASSO oracle (oracle dans le sens que l'estimateur Group

LASSO est estimé sur le plan d'expérience comportant uniquement les variables influentes) a des meilleures performances prédictives que l'estimateur Group LASSO estimé sur un plan d'expérience comportant des faux positifs, et ce même si le plan d'expérience sélectionné par l'estimateur Group LASSO ne comporte que les variables réellement influentes. Belloni *et al.* [16] partent du même constat et combinent, dans le cas des modèles linéaires, le LASSO avec les MCO dans une procédure de type Ante. Comme nous avons pu le constater sur des simulations, Belloni *et al.* [16] constatent une amélioration des performances en prédiction, qui s'explique par le fait que les MCO corrigent le biais introduit par le LASSO.

Nous démontrons ici la consistance en sélection et en estimation de la procédure Post1Bic. Une des difficultés de la démonstration provient du fait que nous travaillons sur une approximation du modèle additif. En effet, nous utilisons des bases de B-Splines pour approcher les composantes additives du modèle et cette approximation introduit un terme d'erreur qu'il faut contrôler.

Après avoir posé quelques notations, nous donnons les hypothèses nécessaires à la démonstration de la consistance de Post1Bic. Certaines hypothèses sont très classiques dans le cadre de la régression et concernent le lissage des composantes, les observations et les erreurs. Nous posons aussi une hypothèse sur le nombre de noeuds, qui ne peut pas croître plus vite que le nombre des observations. Nous faisons cette hypothèse pour contrôler l'erreur due à l'approximation du modèle dans des bases de B-Splines. Nous autorisons le nombre de variables explicatives candidates à croître en fonction du nombre d'observations. Cependant, le nombre de variables explicatives candidates ne peut croître qu'à une vitesse moindre qu'une vitesse exponentielle. Cette hypothèse est nécessaire pour que la pénalisation associée au BIC ne soit pas trop forte. Le nombre de covariables influentes est constant, par souci d'identifiabilité. La norme 2 des composantes influentes est minorée, ce qui est nécessaire pour la consistance en sélection. Contrairement à Huang *et al.* [67], nous ne supposons pas que la norme 2 des composantes influentes est strictement supérieure à une constante strictement positive indépendante du nombre d'observations.

La démonstration de la consistance en sélection se fait en deux étapes. Tout d'abord, nous prouvons que, sous nos hypothèses, le BIC est consistant en sélection lorsque des estimateurs MCO sont utilisés pour estimer le modèle additif. Une difficulté de cette partie de la démonstration résulte du fait qu'il faut contrôler le coût du à l'approximation du modèle additif dans des bases de B-Splines. La seconde étape de la démonstration de la consistance est de montrer que lorsque le coefficient de lissage du Group LASSO a une bonne forme, que nous donnons, le BIC pénalisé est suffisamment proche du BIC associé à l'EMCO, ce qui nous permet, en utilisant la définition de l'estimateur MCO de conclure quant à la consistance en sélection lorsque le BIC est utilisé.

3.4 Démonstration pour la consistance asymptotique en sélection du BIC appliqué à l'EMCO

3.4.1 Cas où $S^* \not\subseteq S$

Démonstration Lemme 3.1 *La démonstration du lemme est en deux parties. D'abord, nous minorons $\|\mathbf{U}_{\bar{S}}(\tilde{\mathbf{h}}_{\bar{S}} - \tilde{\mathbf{h}}_{S^*})\|$, puis nous utilisons un résultat de Huang et al. [67].*

Rappelons que nous sommes dans le cas où $S^ \not\subseteq S$ et que $\bar{S} = S^* \cup S$. Nécessairement, au moins une composante non nulle de $\mathbf{b}_{\bar{S}}^*$ est estimée comme nulle dans $\tilde{\mathbf{h}}_{\bar{S}}$.*

Nous avons donc nécessairement

$$\|\tilde{\mathbf{h}}_{\bar{S}} - \mathbf{b}_{\bar{S}}^*\| \geq \min_{j \in S^*} \|\mathbf{b}_{S^*j}^*\|.$$

De plus, $\sup_{x \in [0,1]} |\sum_{k=1}^{m_n} \mathbf{B}_{j,k}^2(x)|^{1/2} = O(m_n^{1/2})$.

Pour tout $j \in S^*$,

$$\begin{aligned} \|f_j^*\|_2 - \left\| \sum_{k=1}^{m_n} b_{jk}^* \tilde{B}_{j,k} \right\|_2 &\leq \|f_j^* - \sum_{k=1}^{m_n} b_{jk}^* \tilde{B}_{j,k}\|_2 \\ &\leq A_4 m_n^{-\nu} \text{ grâce à Stone [106] et parce que nous travaillons sur un compact.} \end{aligned}$$

Nous avons donc

$$\left\| \sum_{k=1}^{m_n} b_{jk}^* \tilde{B}_{j,k} \right\|_2 \geq \|f_j^*\|_2 - A_4 m_n^{-\nu}.$$

Or

$$\|\mathbf{b}_j^*\|_2 \sim \sqrt{m_n} \left\| \sum_{k=1}^{m_n} b_{jk}^* \tilde{B}_{j,k} \right\|_2,$$

car les B -Splines sont uniformément bornées et évaluées en m_n noeuds.

D'où, en utilisant les propriétés d'approximation des B -Splines, nous pouvons conclure que :

$$\min_{j \in S^*} \|f_j^*\|_2 \leq A_5 \left(\frac{1}{\sqrt{m_n}} \min_{j \in S^*} \|\mathbf{b}_j^*\|_2 + m_n^{-\nu+1/2} \right).$$

D'où

$$\min_{j \in S^*} \|\mathbf{b}_j^*\|_2 \geq A_6 \left(\sqrt{m_n} \min_{j \in S^*} \|f_j^*\|_2 - m_n^{-\nu+1/2} \right).$$

Donc uniformément, pour tout modèle $S \not\supseteq S^*$, nous avons

$$\begin{aligned} \|\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*}\| &\geq \|\tilde{\mathbf{b}}_S - \mathbf{b}_S^*\| - \|\mathbf{b}_S^* - \tilde{\mathbf{b}}_{S^*}\| \\ &\geq A_8 \left(\sqrt{m_n} \min_{j \in S^*} \|f_j^*\| - \frac{m_n}{\sqrt{n}} - m_n^{-\nu+1/2} \right). \end{aligned}$$

Nous utilisons le lemme 3 de Huang et al. [67].

Ici, $m_n = O(n^{\frac{1}{2\nu+1}})$, avec $\nu \geq 2$, d'où $0 < \frac{1}{2\nu+1} < 0.5$. De plus, le cardinal de $\bar{S}$ est inférieur ou égal à $2D$, qui est une constante indépendante de n et d . De plus, avec nos hypothèses 1 et 3, nous respectons les conditions (A3) et (A4) de l'article de Huang et al. [67].

Alors, avec une probabilité tendant vers 1, il existe une constante c_1 telle que

$$\frac{nc_1}{m_n} \leq \rho_{\min}(\mathbf{U}_{\bar{S}}^T \mathbf{U}_{\bar{S}}),$$

où $\rho_{\min}(\mathbf{U}_{\bar{S}}^T \mathbf{U}_{\bar{S}})$ est la valeur propre minimale de $(\mathbf{U}_{\bar{S}}^T \mathbf{U}_{\bar{S}})$.

Nous avons donc

$$\|\mathbf{U}_{\bar{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})\|^2 = (\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})^T \mathbf{U}_{\bar{S}}^T \mathbf{U}_{\bar{S}} (\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*}) \geq A_2 \frac{n}{m_n} v_n^2.$$

□

Démonstration Lemme 3.2 Pour démontrer le Lemme 3.2, nous utilisons d'abord l'inégalité de Cauchy-Schwarz, puis la formule de l'erreur d'estimation de la régression additive et le Lemme 3 de Huang et al. [67].

D'après l'inégalité de Cauchy-Schwarz, nous avons :

$$|2(\mathbf{U}_{\bar{S}} \tilde{\mathbf{b}}_{S^*} - \mathbf{f}^*(\mathbf{X}))^T \mathbf{U}_{\bar{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})| \leq 2\|(\mathbf{U}_{\bar{S}} \tilde{\mathbf{b}}_{S^*} - \mathbf{f}^*(\mathbf{X}))\| \|\mathbf{U}_{\bar{S}}(\tilde{\mathbf{b}}_S - \tilde{\mathbf{b}}_{S^*})\|.$$

Etudions chacun des deux termes :

— Etude de $\|(\mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{f}^*(\mathbf{X}))^T\|$:

$$\begin{aligned} \|(\mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{f}^*(\mathbf{X}))^T\|^2 &\leq \|\mathbf{U}_{\tilde{\mathbf{S}}}(\tilde{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{b}_{\mathbf{S}^*}^*)\|^2 + \sum_{j \in S^*} \|f_j^* - \sum_{k=1}^{m_n} b_{jk}^* \tilde{B}_{j,k}\|^2 \\ &\leq C(m_n + nm_n^{-2\nu} + 2\sqrt{nm_n^{-\nu+1/2}} + m_n^{-2\nu}), \end{aligned}$$

d'où

$$\|(\mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{f}^*(\mathbf{X}))^T\|^2 \leq C(m_n + nm_n^{-2\nu} + o(m_n)),$$

et donc

$$\|(\mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{f}^*(\mathbf{X}))^T\| = O_P((m_n + nm_n^{-2\nu})^{1/2}).$$

— Etude de $\|\mathbf{U}_{\tilde{\mathbf{S}}}(\tilde{\mathbf{b}}_{\mathbf{S}} - \tilde{\mathbf{b}}_{\mathbf{S}^*})\|$:

D'après le Lemme 3 de [67], il existe une constante c_2 telle que

$$\frac{nc_2}{m_n} \geq \rho_{\max}(\mathbf{U}_{\tilde{\mathbf{S}}}^T \mathbf{U}_{\tilde{\mathbf{S}}}),$$

où $\rho_{\max}(\mathbf{U}_{\tilde{\mathbf{S}}}^T \mathbf{U}_{\tilde{\mathbf{S}}})$ est la valeur propre maximale de $(\mathbf{U}_{\tilde{\mathbf{S}}}^T \mathbf{U}_{\tilde{\mathbf{S}}})$.

Nous avons donc

$$\|\mathbf{U}_{\tilde{\mathbf{S}}}(\tilde{\mathbf{b}}_{\mathbf{S}} - \tilde{\mathbf{b}}_{\mathbf{S}^*})\| \leq \sqrt{\frac{nc_2}{m_n}} |v_n|.$$

Nous avons donc

$$|2(\mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{f}^*(\mathbf{X}))^T \mathbf{U}_{\tilde{\mathbf{S}}}(\tilde{\mathbf{b}}_{\mathbf{S}} - \tilde{\mathbf{b}}_{\mathbf{S}^*})| \leq C_4 \sqrt{nm_n^{-2\nu}} \sqrt{\frac{n}{m_n}} |v_n|,$$

avec $C_4 > 0$ constante.

Nous pouvons en conclure que

$$-C_4 \sqrt{nm_n^{-2\nu}} \sqrt{\frac{n}{m_n}} |v_n| \leq 2(\mathbf{U}_{\tilde{\mathbf{S}}} \tilde{\mathbf{b}}_{\mathbf{S}^*} - \mathbf{f}^*(\mathbf{X}))^T \mathbf{U}_{\tilde{\mathbf{S}}}(\tilde{\mathbf{b}}_{\mathbf{S}} - \tilde{\mathbf{b}}_{\mathbf{S}^*}) \leq C_4 \sqrt{nm_n^{-2\nu}} \sqrt{\frac{n}{m_n}} |v_n|.$$

Il nous reste à montrer que ce terme est bien d'un ordre inférieur au troisième terme du membre de droite de l'équation (3.9).

$$\begin{aligned} \frac{\sqrt{nm_n^{-2\nu}} \sqrt{\frac{n}{m_n}} |v_n|}{\frac{n}{m_n} v_n^2} &= \frac{m_n^{-\nu+1/2}}{|v_n|} \\ &= C_5 \frac{n^{\frac{1/2-\nu}{2\nu+1/2}}}{|v_n|}, \end{aligned}$$

avec $C_5 > 0$ constante.

Calculons $n^{-\frac{1/2-\nu}{2\nu+1/2}} |v_n|$, et en notant $C_{10} > 0$ une constante :

$$\begin{aligned} n^{-\frac{1/2-\nu}{2\nu+1/2}} |v_n| &\geq C_{10} \left(n^{-\frac{1/2-\nu}{2\nu+1}} n^{1/2 \frac{1}{2\nu+1}} \min_{j \in S^*} \|f_j^*\|_2 - 2n^{(-\nu+1/2)/(2\nu+1)} n^{-\frac{1/2-\nu}{2\nu+1}} \right) \\ &>> C_{10} \left(\sqrt{\log(nd)} - 2n^{\frac{-\nu}{2\nu+1}} \right). \end{aligned}$$

Or $\sqrt{\log(nd)} \xrightarrow{n \rightarrow +\infty} \infty$ et $n^{\frac{-\nu}{2\nu+1}} \xrightarrow{n \rightarrow +\infty} 0$. Donc

$$n^{-\frac{1/2-\nu}{2\nu+1/2}} |v_n| >> C_{10} \left(\sqrt{\log(nd)} - 2n^{\frac{-\nu}{2\nu+1}} \right) \xrightarrow{n \rightarrow +\infty} +\infty,$$

et donc

$$\frac{\sqrt{nm_n^{-2\nu}} \sqrt{\frac{n}{m_n}} |v_n|}{\frac{n}{m_n} v_n^2} \xrightarrow{n \rightarrow +\infty} 0.$$

D'où asymptotiquement, le second terme de l'équation (3.9) est négligeable devant le troisième terme de (3.9).

□

Proposition 3.1 Pour tout $\bar{S}$ tel que $S^* \subseteq \bar{S}$, nous avons

$$O_p(\|\Pi_{\bar{S}} B\|^2) = O_p(\|B\|^2) = O_p(nm_n^{-2\nu}). \quad (3.16)$$

Proposition 3.2 Soit M constante telle que $M < n$ et $\bar{S}$ tel que $S^* \subseteq \bar{S}$, alors

$$\sup_{\bar{S} \mid \text{card}(\bar{S}) < M} \underline{\epsilon}^T \Pi_{\bar{S}} \underline{\epsilon} = o_P(m_n \log(nd)). \quad (3.17)$$

De plus, $m_n \log(nd)$ est d'un ordre plus petit que $\frac{n}{m_n} v_n^2$.

Les deux propositions précédentes permettent de montrer que, lorsque le vrai plan d'expérience S^* est inclus dans $\bar{S}$, alors $\|\underline{Y} - \underline{U}_{\bar{S}} \hat{\underline{b}}_{\bar{S}}\|^2/n - \|\underline{Y} - \underline{U}_{\bar{S}} \underline{b}_{\bar{S}}^*\|^2/n$ tend vers 0 avec une probabilité tendant vers 1, c'est-à-dire :

Proposition 3.3 Soit $\bar{S}$ tel que $S^* \subseteq \bar{S}$,

alors,

$$\|\underline{Y} - \underline{U}_{\bar{S}} \hat{\underline{b}}_{\bar{S}}\|^2 - \|\underline{Y} - \underline{U}_{\bar{S}} \underline{b}_{\bar{S}}^*\|^2 = O_p(nm_n^{-2\nu}) + o_p(m_n \log(nd)). \quad (3.18)$$

Montrons les trois propositions dans l'ordre.

Démonstration Proposition 3.1 Nous avons $\Pi_{\bar{S}} \underline{U}_{\bar{S}} = \underline{U}_{\bar{S}}, \Pi_{\bar{S}}^T \Pi_{\bar{S}} = \Pi_{\bar{S}}^T, (\Pi_{\bar{S}}^T \Pi_{\bar{S}})^T = \Pi_{\bar{S}}^T \Pi_{\bar{S}}$ et $(\Pi_{\bar{S}}^T \Pi_{\bar{S}})^T = (\Pi_{\bar{S}}^T)^T = \Pi_{\bar{S}}$. Nous avons donc $\Pi_{\bar{S}} = (\Pi_{\bar{S}}^T \Pi_{\bar{S}})^T = \Pi_{\bar{S}}^T \Pi_{\bar{S}} = \Pi_{\bar{S}}^T$. Enfin, $(I - \Pi_{\bar{S}})$ est un projecteur.

$$\begin{aligned} \|\Pi_{\bar{S}} B\|^2 &= B^T \Pi_{\bar{S}}^T \Pi_{\bar{S}} B \\ &= (\underline{f}^*(\underline{X}) - \underline{U}_{\bar{S}} \underline{b}_{\bar{S}}^*)^T \Pi_{\bar{S}} (\underline{f}^*(\underline{X}) - \underline{U}_{\bar{S}} \underline{b}_{\bar{S}}^*) \\ &= \underline{f}^*(\underline{X})^T (\Pi_{\bar{S}} - I) \underline{f}^*(\underline{X}) + \underline{f}^*(\underline{X})^T \underline{f}^*(\underline{X}) - \underline{b}_{\bar{S}}^{*T} \underline{U}_{\bar{S}}^T \underline{f}^*(\underline{X}) - \underline{f}^*(\underline{X})^T \underline{U}_{\bar{S}} \underline{b}_{\bar{S}}^* + \underline{b}_{\bar{S}}^{*T} \underline{U}_{\bar{S}}^T \underline{U}_{\bar{S}} \underline{b}_{\bar{S}}^* \\ &= -\underline{f}^*(\underline{X})^T (I - \Pi_{\bar{S}}) \underline{f}^*(\underline{X}) + \|B\|^2 \\ &= -\underline{f}^*(\underline{X})^T (I - \Pi_{\bar{S}})^T (I - \Pi_{\bar{S}}) \underline{f}^*(\underline{X}) + \|B\|^2 \\ &= \|B\|^2 - \|(I - \Pi_{\bar{S}}) \underline{f}^*(\underline{X})\|^2 \\ &\leq \|B\|^2. \end{aligned}$$

D'où

$$\|\Pi_{\bar{S}} B\|^2 = O_p(\|B\|^2).$$

De plus,

$$\begin{aligned}
\|\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) - \mathbf{U}_{\bar{\mathcal{S}}} \underline{\mathbf{b}}_{\bar{\mathcal{S}}}^*\|^2 &= \left\| \sum_{j \in \bar{\mathcal{S}}} \underline{\mathbf{f}}_j^*(\underline{\mathbf{X}}^j) - \sum_{j \in \bar{\mathcal{S}}} \sum_{k=1}^{m_n} b_{j,k} (B_{jk}(x_{ij}) - \bar{B}_{jk}) \right\|^2 \\
&\leq \sum_{j \in \bar{\mathcal{S}}} \|\underline{\mathbf{f}}_j^*(\underline{\mathbf{X}}^j) - \sum_{k=1}^{m_n} b_{j,k} (B_{jk} - \bar{B}_{jk})\|^2 \\
&\leq \sum_{j \in \bar{\mathcal{S}}} n \|f_j^* - \sum_{k=1}^{m_n} b_{j,k} (B_{jk} - \bar{B}_{jk})\|_{+\infty}^2 \\
&\leq \sum_{j \in \bar{\mathcal{S}}} nm_n^{-2\nu} \\
&\leq D_1 nm_n^{-2\nu} \text{ avec } D_1 \text{ constante car } s^* \text{ supposé constant.}
\end{aligned}$$

□

Démonstration Proposition 3.2 Nous avons supposé que le vecteur des erreurs était gaussien centré, indépendant et avec des composantes de variance bornée par σ^2 p.s. (voir l'hypothèse 2). $\|\Pi_{\bar{\mathcal{S}}} \underline{\epsilon}\|^2 = \underline{\epsilon}^T \Pi_{\bar{\mathcal{S}}} \underline{\epsilon}$ est donc une forme quadratique d'un vecteur gaussien. La dimension de l'espace propre de $\Pi_{\bar{\mathcal{S}}}$ est au plus M avec M fini. Cette forme quadratique peut donc s'exprimer comme une somme de variables du Chi-deux indépendantes.

Calculons l'espérance de $\underline{\epsilon}^T \Pi_{\bar{\mathcal{S}}} \underline{\epsilon}$.

$$\begin{aligned}
E(\underline{\epsilon}^T \Pi_{\bar{\mathcal{S}}} \underline{\epsilon}) &= E\left(\sum_{i=1}^n \sum_{j=1}^n (\Pi_{\bar{\mathcal{S}}})_{ij} \epsilon_j \epsilon_i\right) \\
&\leq M\sigma^2.
\end{aligned}$$

Grâce à l'inégalité de Markov, nous pouvons en conclure que

$$P\left(\sup_{\bar{\mathcal{S}} | \text{card}(\bar{\mathcal{S}}) \leq M} \underline{\epsilon}^T \Pi_{\bar{\mathcal{S}}} \underline{\epsilon} \geq m_n \log(nd)\right) \leq \frac{M\sigma^2}{m_n \log(nd)} \xrightarrow{n \rightarrow +\infty} 0.$$

Montrons que $m_n \log(nd)$ est d'un ordre plus petit ou égal que $\frac{n}{m_n} v_n^2$

$$\frac{m_n \log(nd)}{\frac{n}{m_n} \left(\sqrt{m_n} \min \|f_j^*\|_2 - 2m_n^{-\nu+1/2}\right)^2} = \frac{m_n}{n} \frac{\log(nd)}{(\min \|f_j^*\|_2)^2} \frac{1}{\left(1 - 2\frac{m_n^{-\nu}}{\min \|f_j^*\|_2}\right)^2},$$

or

$$\frac{m_n^{-\nu}}{\min \|f_j^*\|_2} < G_1 \frac{m_n^{-\nu}}{\sqrt{m_n^{-2\nu} \log(nd)}} \xrightarrow{n \rightarrow +\infty} 0 \text{ d'après l'hypothèse 5 et } G_1 > 0 \text{ constante.}$$

D'où

$$\frac{1}{\left(1 - 2\frac{m_n^{-\nu}}{\min \|f_j^*\|_2}\right)^2} \xrightarrow{n \rightarrow +\infty} 1.$$

De plus, $\frac{\log(nd)}{(\min \|f_j^*\|_2)^2} < \frac{\log(nd)}{m_n^{-2\nu} \log(nd)} < m_n^{2\nu}$.

Finalement,

$$\frac{m_n \log(nd)}{\frac{n}{m_n} \left(\sqrt{m_n} \min \|f_j^*\|_2 - 2m_n^{-\nu+1/2} \right)^2} < \frac{m_n^{2\nu+1}}{n} = A.$$

D'où $m_n \log(nd) = O(\frac{n}{m_n} v_n^2)$.

D'où $\sup_{\bar{S} | \text{card}(\bar{S})} \underline{\epsilon}^T \Pi_{\bar{S}} \underline{\epsilon} = o_P(\frac{n}{m_n} v_n^2)$.

□

Démonstration Proposition 3.3

$$\begin{aligned} \|\underline{Y} - \mathbf{U}_{\bar{S}} \hat{\mathbf{b}}_{\bar{S}}\|^2 &= \|\underline{Y} - \mathbf{U}_{\bar{S}} (\mathbf{U}_{\bar{S}}^T \mathbf{U}_{\bar{S}})^{-1} \mathbf{U}_{\bar{S}}^T \underline{Y}\|^2 \\ &= \|(I - \Pi_{\bar{S}}) \underline{Y}\|^2, \end{aligned}$$

et

$$\|\underline{Y} - \mathbf{U}_{\bar{S}} \mathbf{b}_{\bar{S}}^*\|^2 = \|B + \underline{\epsilon}\|^2,$$

d'où

$$\begin{aligned} \|\underline{Y} - \mathbf{U}_{\bar{S}} \hat{\mathbf{b}}_{\bar{S}}\|^2 - \|\underline{Y} - \mathbf{U}_{\bar{S}} \mathbf{b}_{\bar{S}}^*\|^2 &= \|(I - \Pi_{\bar{S}}) \underline{Y}\|^2 - \|B + \underline{\epsilon}\|^2 \\ &= \|(I - \Pi_{\bar{S}})(\mathbf{U}_{\bar{S}} \mathbf{b}_{\bar{S}}^* + B + \underline{\epsilon})\|^2 - \|B + \underline{\epsilon}\|^2 \\ &= \|(I - \Pi_{\bar{S}})(B + \underline{\epsilon})\|^2 - \|B + \underline{\epsilon}\|^2 \\ &= (B + \underline{\epsilon})^T (I - \Pi_{\bar{S}})^T (I - \Pi_{\bar{S}}) (B + \underline{\epsilon}) - (B + \underline{\epsilon})^T (B + \underline{\epsilon}) \\ &= -(B + \underline{\epsilon})^T \Pi_{\bar{S}} (B + \underline{\epsilon}) \text{ (car } \Pi_{\bar{S}}^T \Pi_{\bar{S}} = \Pi_{\bar{S}} \text{ et } \Pi_{\bar{S}}^T = \Pi_{\bar{S}}) \\ &= -\|\Pi_{\bar{S}}(B + \underline{\epsilon})\|^2. \end{aligned}$$

Comme $\|X + Y\|^2 \leq 2\|X\|^2 + 2\|Y\|^2$ et $\|X + Y\|^2 \geq 2\|X\|^2 - 2\|Y\|^2$, nous avons donc

$$\|\underline{Y} - \mathbf{U}_{\bar{S}} \hat{\mathbf{b}}_{\bar{S}}\|^2 - \|\underline{Y} - \mathbf{U}_{\bar{S}} \mathbf{b}_{\bar{S}}^*\|^2 \geq -2(\|\Pi_{\bar{S}} B\|^2 + \|\Pi_{\bar{S}} \underline{\epsilon}\|^2),$$

et

$$\|\underline{Y} - \mathbf{U}_{\bar{S}} \hat{\mathbf{b}}_{\bar{S}}\|^2 - \|\underline{Y} - \mathbf{U}_{\bar{S}} \mathbf{b}_{\bar{S}}^*\|^2 \leq 2(\|\Pi_{\bar{S}} B\|^2 + \|\Pi_{\bar{S}} \underline{\epsilon}\|^2).$$

Comme $S^* \subseteq \bar{S}$, d'après les Propositions 3.1 et 3.2, nous avons

$$O_P(\|\Pi_{\bar{S}} B\|^2) = O_P(\|B\|^2) = O_P(nm_n^{-2\nu}) \text{ et } \sup_{\bar{S} | \text{card}(\bar{S}) < M} \underline{\epsilon}^T \Pi_{\bar{S}} \underline{\epsilon} = o_P(m_n \log(nd)).$$

□

Démonstration Lemme 3.3 Puisque pour tout couple de réels (a, b) , $-4a^2 - \frac{1}{4}b^2 \leq -2ab$, nous avons,

$$-2\underline{\epsilon}^T \mathbf{U}_{\bar{S}} (\tilde{\mathbf{b}}_{\bar{S}} - \tilde{\mathbf{b}}_{S^*}) \geq -4\underline{\epsilon}^T \Pi_{\bar{S}} \Pi_{\bar{S}} \underline{\epsilon} - \frac{1}{4} \|\mathbf{U}_{\bar{S}} (\tilde{\mathbf{b}}_{\bar{S}} - \tilde{\mathbf{b}}_{S^*})\|^2.$$

Donc

$$\begin{aligned} -2\underline{\epsilon}^T \mathbf{U}_{\bar{S}} (\tilde{\mathbf{b}}_{\bar{S}} - \tilde{\mathbf{b}}_{S^*}) + \|\mathbf{U}_{\bar{S}} (\tilde{\mathbf{b}}_{\bar{S}} - \tilde{\mathbf{b}}_{S^*})\|^2 &\geq -4\underline{\epsilon}^T \Pi_{\bar{S}} \Pi_{\bar{S}} \underline{\epsilon} + \frac{3}{4} \|\mathbf{U}_{\bar{S}} (\tilde{\mathbf{b}}_{\bar{S}} - \tilde{\mathbf{b}}_{S^*})\|^2 \\ &\geq \frac{3}{4} \frac{n}{m_n} v_n^2 + o_P\left(\frac{n}{m_n} v_n^2\right) \quad \text{d'après le Lemme 3.2.} \end{aligned}$$

Ceci démontre le Lemme 3.3.

□

Démonstration Lemme 3.4 *Nous commençons par montrer le résultat pour*

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \mathbf{b}_{\mathbf{S}^*}\|^2/n.$$

Nous avons

$$\begin{aligned} \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \mathbf{b}_{\mathbf{S}^*}\|^2/n &= \|\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) + \underline{\boldsymbol{\epsilon}} - \mathbf{U}_{\mathbf{S}^*} \mathbf{b}_{\mathbf{S}^*}\|^2/n \quad \text{par définition de } \underline{\mathbf{Y}} \\ &\geq \frac{\|\underline{\boldsymbol{\epsilon}}\|^2}{n} - \|\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) - \mathbf{U}_{\mathbf{S}^*} \mathbf{b}_{\mathbf{S}^*}\|^2/n \\ &\geq \frac{\|\underline{\boldsymbol{\epsilon}}\|^2}{n} + O_P(m_n^{-2\nu}). \end{aligned}$$

Nous allons montrer que $\frac{\|\underline{\boldsymbol{\epsilon}}\|^2}{n}$ est strictement supérieur à 0 avec une probabilité tendant vers 1. Pour cela, nous allons utiliser l'inégalité de Bienaymé-Tchebychev :

Soit X une variable aléatoire de variance finie, alors pour tout $\alpha > 0$,

$$P(|X - E(X)| \geq \alpha) \leq \frac{V(X)}{\alpha}.$$

Nous considérons la variable aléatoire $\frac{1}{n} \sum \epsilon_i^2$. D'après l'hypothèse 2,

$$E\left(\frac{1}{n} \sum \epsilon_i^2\right) = \frac{1}{n} \sum \sigma^2(\underline{\mathbf{X}}_i),$$

et, par indépendance et comme ϵ_i suit une loi normale centrée de variance $\sigma^2(\underline{\mathbf{X}}_i)$, le moment d'ordre 4 de ϵ_i est $3\sigma^4(\underline{\mathbf{X}}_i)$.

D'où

$$\begin{aligned} V\left(\frac{1}{n} \sum \epsilon_i^2\right) &= \frac{3}{n^2} \sum \sigma^4(\underline{\mathbf{X}}_i) \\ &\leq \frac{3}{n} \sigma^4 \quad \text{d'après l'hypothèse 2.} \end{aligned}$$

Finalement, d'après Bienaymé-Tchebychev, pour tout $\alpha > 0$,

$$P\left(\left|\frac{1}{n} \sum \epsilon_i^2 - \frac{1}{n} \sum \sigma^2(\underline{\mathbf{X}}_i)\right| \geq \alpha\right) \leq \frac{3}{n} \frac{\sigma^4}{\alpha} \rightarrow 0,$$

d'où

$$\frac{1}{n} \sum \epsilon_i^2 \xrightarrow{P} \frac{1}{n} \sum \sigma^2(\underline{\mathbf{X}}_i) \geq \frac{n}{n} \sigma_1^2 \geq \sigma_1^2 > 0 \quad \text{d'après l'hypothèse 2.}$$

Donc, avec une probabilité tendant vers 1, $\|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^} \mathbf{b}_{\mathbf{S}^*}\|^2/n$ est bornée inférieurement par une constante strictement positive.*

Nous passons à l'étude de la suite $\inf_{\mathbf{S}^ \subseteq \mathbf{S} \mid \text{card}(\mathbf{S}) \leq D} \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \hat{\mathbf{b}}_{\mathbf{S}}\|^2/n$.*

$\mathbf{S}^ \subseteq \mathbf{S}$, d'où, d'après le Lemme 3.3,*

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \hat{\mathbf{b}}_{\mathbf{S}}\|^2/n = \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \mathbf{b}_{\mathbf{S}}^*\|^2/n + O_P(m_n^{-2\nu}) + o_P\left(\frac{m_n}{n} \log(nd)\right) = \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \mathbf{b}_{\mathbf{S}}^*\|^2/n + o_P(1).$$

Or, par construction,

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}} \mathbf{b}_{\mathbf{S}}^*\|^2/n = \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \mathbf{b}_{\mathbf{S}^*}\|^2/n \geq \sigma_1^2 + O_P(m_n^{-2\nu}),$$

d'où

$$\|\underline{\mathbf{Y}} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2 / n \geq \sigma_1^2 + o_P(1) > 0,$$

avec une probabilité tendant vers 1 dès que $S^* \subseteq S$.

D'où la conclusion du lemme.

□

Démonstration Théorème 3.1 *Nous avons*

$$BIC(S) - BIC(S^*) \geq \min\left(\frac{L_1}{m_n} v_n^2, \log(2)\right) + (\text{card}(S) - \text{card}(S^*)) \frac{m_n \log(nd)}{n},$$

où $L_1 > 0$ constante.

Or, lorsque $n \rightarrow +\infty$, $L_1 m_n v_n^2 > 0$ et

$$\frac{m_n \log(nd)}{n} = A \frac{n^{1/(2\nu+1)}}{n} \log(nd) = A n^{-2\nu/(2\nu+1)} (\log(n) + \log(d)).$$

Or, d'après les hypothèses 5, $n^{-2\nu/(2\nu+1)} \log(d) \xrightarrow{n \rightarrow +\infty} 0$ et il est évident que $n^{-2\nu/(2\nu+1)} \log(n) \xrightarrow{n \rightarrow +\infty} 0$.

De plus, $|\text{card}(S) - \text{card}(S^*)| \leq |M - D|$, or M et D indépendantes de n et finies, d'où

$$|\text{card}(S) - \text{card}(S^*)| \frac{m_n \log(nd)}{n} \xrightarrow{n \rightarrow +\infty} 0.$$

Et de plus,

$$|\text{card}(S) - \text{card}(S^*)| \frac{m_n \log(nd)}{n} = o(\min \|f_j^*\|^2).$$

D'où, avec probabilité 1, lorsque $n \rightarrow +\infty$, pour tout S tel que $S^* \not\subseteq S$,

$$BIC(S) - BIC(S^*) > 0.$$

D'où la conclusion du théorème.

□

3.4.2 Cas où $S^* \subseteq S$

Démonstration Lemme 3.5 *Commençons par l'étude de*

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2.$$

Tout d'abord, nous constatons que $\mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^* = \mathbf{U}_S \underline{\mathbf{b}}_S^*$, et donc

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*\|^2 = \|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S^*\|^2.$$

De toute évidence, nous avons bien

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2 = \|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \hat{\mathbf{b}}_S\|^2 \geq 0,$$

par définition de l'estimateur MCO, et dans nos conditions, le programme des MCO est identifiable. Nous avons égalité si, et seulement si, $\hat{\mathbf{b}}_S = \underline{\mathbf{b}}_S^*$.

Or $S^* \subseteq S$, donc d'après la Proposition 3.3,

$$\|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2 = O_P(nm_n^{-2\nu}) + o_P(m_n \log(nd)).$$

Comme $S^* \subseteq S$, nous pouvons appliquer le Lemme 3.4. Avec une probabilité tendant vers 1, $\|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2/n$ tend vers une constante strictement supérieure à 0. Or

$$O_P(m_n^{-2\nu}) + o_P(m_n \log(nd)/n) \xrightarrow{P} 0.$$

Donc, avec une probabilité tendant vers 1,

$$0 \leq \frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2}{\|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2} \xrightarrow{P} 0.$$

□

Démonstration Théorème 3.2 D'après le Lemme 3.5, avec une probabilité tendant vers 1,

$$\frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2}{\|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2} \xrightarrow{P} 0,$$

et

$$\frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2}{\|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2} = O_P(m_n^{-2\nu}) + o_P(m_n \log(nd)/n).$$

Or, si $x_n \rightarrow 0$, $\log(1 + x_n) = x_n + o(x_n)$ pour n suffisamment grand.

Comme

$$BIC(S^*) - BIC(S) \leq \log \left(1 + \frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2}{\|\underline{\mathbf{Y}} - \mathbf{U}_S \underline{\mathbf{b}}_S\|^2} \right) + (\text{card}(S^*) - \text{card}(S)) m_n \frac{\log(nd)}{n},$$

nous avons

$$\begin{aligned} BIC(S^*) - BIC(S) &\leq O_P(m_n^{-2\nu}) + o_P\left(\frac{m_n}{n} \log(nd)\right) \\ &\quad + o(O_P(m_n^{-2\nu}) + o_P(m_n \log(nd)/n)) \\ &\quad + (\text{card}(S^*) - \text{card}(S)) m_n \frac{\log(nd)}{n}. \end{aligned}$$

Or si $S^* \subsetneq S$, alors

$$(\text{card}(S^*) - \text{card}(S)) < 0,$$

car les cardinaux sont finis par hypothèse.

Montrons que $O_P(m_n^{-2\nu})$ est négligeable par rapport à $m_n \frac{\log(nd)}{n}$.

$$\begin{aligned} \frac{O_P(m_n^{-2\nu})}{m_n \frac{\log(nd)}{n}} &= \frac{O_P(n^{-2\nu/(2\nu+1)})}{n^{1/(2\nu+1)} \frac{\log(nd)}{n}} \\ &= \frac{O_P(1)}{\log(nd)} \\ &\xrightarrow[n \rightarrow +\infty]{P} 0, \end{aligned}$$

donc

$$BIC(S^*) - BIC(S) < 0,$$

avec une probabilité tendant vers 1.

□

Démonstration Théorème 3.3 D'après les Théorèmes 3.1 et 3.2, le BIC est minimum lorsque l'estimateur MCO est calculé sur S^* avec une probabilité tendant vers 1, lorsque le cardinal maximal des plans d'expérience candidats est inférieur à une constante indépendante de n .

□

3.5 Démonstration de la consistance en sélection de la procédure Post1

Proposition 3.4 Notons

$$\Omega_0 = \{\|\hat{\Sigma}_j^{1/2} - I_m\| \leq 0.5, 1 \leq j \leq d\}.$$

Supposons que

$$\frac{m_n \log(d)}{n} \rightarrow 0.$$

Alors $P(\Omega_0) \rightarrow 1$.

Et de plus,

$$\max_{1 \leq j \leq d} \|\hat{\Sigma}_j^{1/2} - I_m\| \xrightarrow{P} 0,$$

et

$$\max_{1 \leq j \leq d} \|\bar{\mathbf{B}}_j\| \xrightarrow{P} 0.$$

Démonstration Proposition 3.4 Le résultat découle de Kato 2011 [71] et 2012 [72] (Lemme B.1, page 40).

□

Proposition 3.5 Si

$$\tilde{\mathbf{b}}_\lambda = \arg \min_{\mathbf{b} \in \mathbf{R}^{dm_n+1}} \|\mathbf{Y} - \mathbf{U}^T \mathbf{b}\|^2 + \sqrt{m_n} \lambda \sum_{j=1}^d \|\mathbf{b}_j\|_2,$$

et $\text{card}(S_\lambda) \leq M$,

alors

$$\hat{\mathbf{b}}_\lambda = \arg \min_{\mathbf{b} \in \mathbf{R}^{\text{card}(S_\lambda)m_n+1}} \|\mathbf{Y} - \mathbf{U}_{S_\lambda}^T \mathbf{b}\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\mathbf{b}_j\|_2.$$

Démonstration Proposition 3.5 Pour démontrer ce résultat, nous montrons qu'il n'existe pas de $\tilde{\mathbf{a}}_\lambda$ tel que

$$\|\mathbf{Y} - \mathbf{U}_{S_\lambda}^T \hat{\mathbf{b}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\hat{\mathbf{b}}_{\lambda,j}\|_2 > \|\mathbf{Y} - \mathbf{U}_{S_\lambda}^T \tilde{\mathbf{a}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\tilde{\mathbf{a}}_{\lambda,j}\|_2.$$

Nous faisons un raisonnement par l'absurde. Supposons qu'il existe $\tilde{\mathbf{a}}_\lambda$ tel que

$$\|\mathbf{Y} - \mathbf{U}_{S_\lambda}^T \hat{\mathbf{b}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\hat{\mathbf{b}}_{\lambda,j}\|_2 > \|\mathbf{Y} - \mathbf{U}_{S_\lambda}^T \tilde{\mathbf{a}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\tilde{\mathbf{a}}_{\lambda,j}\|_2.$$

Notons $\bar{\mathbf{a}}_\lambda$ tel que $\bar{\mathbf{a}}_{\lambda,j} = \tilde{\mathbf{a}}_{\lambda,j}$ si $j \in S_\lambda$, $= \mathbf{0}$ sinon.

Alors il est clair que

$$\sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\tilde{\mathbf{a}}_{\lambda,j}\|_2 = \sqrt{m_n} \lambda \sum_{j=1}^d \|\bar{\mathbf{a}}_{\lambda,j}\|_2,$$

par définition de $\tilde{\mathbf{a}}_\lambda$.

De même, puisque $\tilde{\mathbf{a}}_{\lambda, \mathbf{j}} = \tilde{\mathbf{a}}_{\lambda, \mathbf{j}}$ si $j \in S_\lambda$, nous avons

$$\|\underline{\mathbf{Y}} - \mathbf{U}\tilde{\mathbf{a}}_\lambda\|^2 = \|\underline{\mathbf{Y}} - \mathbf{U}_{S_\lambda}^T \tilde{\mathbf{a}}_\lambda\|^2.$$

De même,

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{S_\lambda}^T \hat{\mathbf{b}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\hat{\mathbf{b}}_{\lambda, \mathbf{j}}\|_2 = \|\underline{\mathbf{Y}} - \mathbf{U}\tilde{\mathbf{b}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j=1}^d \|\tilde{\mathbf{b}}_{\lambda, \mathbf{j}}\|_2,$$

puisque

$$\|\underline{\mathbf{Y}} - \mathbf{U}_{S_\lambda}^T \hat{\mathbf{b}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\hat{\mathbf{b}}_{\lambda, \mathbf{j}}\|_2 > \|\underline{\mathbf{Y}} - \mathbf{U}_{S_\lambda}^T \tilde{\mathbf{a}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j \in S_\lambda} \|\tilde{\mathbf{a}}_{\lambda, \mathbf{j}}\|_2,$$

nous avons donc

$$\|\underline{\mathbf{Y}} - \mathbf{U}\tilde{\mathbf{b}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j=1}^d \|\tilde{\mathbf{b}}_{\lambda, \mathbf{j}}\|_2 > \|\underline{\mathbf{Y}} - \mathbf{U}\tilde{\mathbf{a}}_\lambda\|^2 + \sqrt{m_n} \lambda \sum_{j=1}^d \|\tilde{\mathbf{a}}_{\lambda, \mathbf{j}}\|_2,$$

ce qui est contradictoire avec

$$\tilde{\mathbf{b}}_\lambda = \arg \min_{\mathbf{b} \in \mathbf{R}^{dm_n+1}} \|\underline{\mathbf{Y}} - \mathbf{U}^T \mathbf{b}\|^2 + \sqrt{m_n} \lambda \sum_{j=1}^d \|\mathbf{b}_j\|_2.$$

Ceci permet de conclure la démonstration de la proposition.

□

Proposition 3.6 $\tilde{\mathbf{b}}_{\lambda_n}$ est solution de (3.15) si

$$\begin{cases} \mathbf{U}_j^T (\underline{\mathbf{Y}} - \mathbf{U}\tilde{\mathbf{b}}_{\lambda_n}) = \lambda_n \sqrt{m_n} \frac{\tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}}{\|\tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\|} \text{ si } j \in S_{\lambda_n}, \\ \|\mathbf{U}_j^T (\underline{\mathbf{Y}} - \mathbf{U}\tilde{\mathbf{b}}_{\lambda_n})\| \leq \lambda_n \sqrt{m_n} \text{ si } j \notin S_{\lambda_n}. \end{cases} \quad (3.19)$$

Soit $w_n = (\tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}} / \|\tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\|, j \in S^*)$. Notons

$$\hat{\mathbf{b}}_{\lambda_n, S^*} = (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} (\mathbf{U}_{S^*}^T \underline{\mathbf{Y}} - \lambda_n \sqrt{m_n} w_n).$$

Si $S_{\lambda_n} = S^*$, alors d'après le Lemme 3.5, si S^* est composé des s^* premières covariables, $\tilde{\mathbf{b}}_{\lambda_n} = (\hat{\mathbf{b}}_{\lambda_n, S^*}^T, \mathbf{0}^T)^T$.

Comme $\mathbf{U}\tilde{\mathbf{b}}_{\lambda_n} = \mathbf{U}_{S^*} \hat{\mathbf{b}}_{\lambda_n, S^*}$ et $\{\mathbf{U}_j, j \in S^*\}$ sont linéairement indépendants,

$S_{\lambda_n} = S^*$ si

$$\begin{cases} \|\mathbf{b}_j^*\| - \|\tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\| < \|\mathbf{b}_j^*\| \text{ si } j \in S^*, \\ \|\mathbf{U}_j^T (\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{\lambda_n, S^*})\| \leq \lambda_n \sqrt{m_n} \text{ si } j \notin S_{\lambda_n}. \end{cases}$$

Démonstration Proposition 3.6 La première partie du lemme est classique (voir la Proposition 8 de Bach [13] par exemple) et consiste en l'écriture des équations normales du Group LASSO.

$S_{\lambda_n} = S^*$ si

$$\begin{cases} \|\mathbf{b}_j^*\| - \|\tilde{\mathbf{b}}_{\lambda_n, \mathbf{j}}\| < \|\mathbf{b}_j^*\| \text{ si } j \in S^*, \\ \|\mathbf{U}_j^T (\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{\lambda_n, S^*})\| \leq \lambda_n \sqrt{m_n} \text{ si } j \notin S_{\lambda_n}. \end{cases}$$

ce qui est démontré au début de la démonstration par Huang et al. [67].

□

Démonstration Lemme 3.6 Soit $j \in S^*$.

Notons

$$T_j = (\mathbf{0}_{m_n \times m_n}, \dots, \mathbf{0}_{m_n \times m_n}, \mathbf{I}_{m_n \times m_n}, \mathbf{0}_{m_n \times m_n}, \dots, \mathbf{0}_{m_n \times m_n}),$$

avec $\mathbf{0}_{m_n \times m_n}$ matrice nulle de taille $m_n \times m_n$ et $\mathbf{I}_{m_n \times m_n}$ matrice identité de taille $m_n \times m_n$. Le jème bloc de T_j est $\mathbf{I}_{m_n \times m_n}$.

$$\begin{aligned} \tilde{\mathbf{b}}_{\lambda_n \mathbf{j}} &= T_j \tilde{\mathbf{b}}_{\lambda_n} \\ &= T_j (\hat{\mathbf{b}}_{S^*} - (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}). \end{aligned}$$

Notons $C_{S^*} = \frac{1}{n} \mathbf{U}_{S^*}^T \mathbf{U}_{S^*}$, alors $\|\tilde{\mathbf{b}}_{\lambda_n \mathbf{j}} - \underline{\mathbf{b}}_{S^* \mathbf{j}}^*\| = \|T_j (\hat{\mathbf{b}}_{S^*} - \frac{1}{n} C_{S^*}^{-1} \underline{\mathbf{r}} - \underline{\mathbf{b}}_{S^*}^*)\| \leq \|T_j (\hat{\mathbf{b}}_{S^*} - \underline{\mathbf{b}}_{S^*}^*)\| + \frac{1}{n} \|T_j C_{S^*}^{-1} \underline{\mathbf{r}}\|$, or $\|\hat{\mathbf{b}}_{S^*} - \underline{\mathbf{b}}_{S^*}^*\| = O(\frac{m_n}{\sqrt{n}})$, d'après Stone [106], d'où

$$\max_j \|T_j (\hat{\mathbf{b}}_{S^*} - \underline{\mathbf{b}}_{S^*}^*)\| \leq \frac{m_n}{\sqrt{n}}.$$

De plus,

$$\begin{aligned} \frac{1}{n} \|T_j C_{S^*}^{-1} \underline{\mathbf{r}}\| &\leq \frac{1}{n} \|T_j\| \|C_{S^*}^{-1}\| \|\underline{\mathbf{r}}\| \\ &\leq O(1) \frac{1}{n} m_n \lambda_n \sqrt{m_n} \\ &\leq O(1) \frac{m_n}{\sqrt{n}} \sqrt{\log(m_n^2 d)} \\ &\leq O(1) \frac{m_n}{\sqrt{n}} \sqrt{\log(nd)} \\ &\leq o(\sqrt{m_n} \min \|f_j^*\|), \end{aligned}$$

d'où

$$\|\tilde{\mathbf{b}}_{\lambda_n \mathbf{j}} - \underline{\mathbf{b}}_{S^* \mathbf{j}}^*\| \leq \frac{m_n}{\sqrt{n}} + o(\sqrt{m_n} \min \|f_j^*\|).$$

De plus,

$$\begin{aligned} \min_{j \in S^*} \|\mathbf{b}_j^*\| &\geq A_6 \left(\sqrt{m_n} \min_{j \in S^*} \|f_j^*\| - O(1) \frac{m_n}{\sqrt{n}} \right) \\ &\geq O(1) \left(\sqrt{m_n} \sqrt{\frac{m_n}{n} \log(nd)} + O\left(\frac{m_n}{\sqrt{n}}\right) \right) \\ &\geq O(1) \left(\frac{m_n}{\sqrt{n}} \sqrt{\log(nd)} + O\left(\frac{m_n}{\sqrt{n}}\right) \right). \end{aligned}$$

Or, il est clair que

$$\frac{m_n}{\sqrt{n}} = o\left(\frac{m_n}{\sqrt{n}} \sqrt{\log(nd)}\right),$$

et donc

$$P(\text{il existe } j \in S^*, \|\underline{\mathbf{b}}_j^* - \tilde{\mathbf{b}}_{\lambda_n \mathbf{j}}\| \geq \|\underline{\mathbf{b}}_j^*\|) \rightarrow 0.$$

□

Pour démontrer le résultat suivant, nous supposons que pour tout $j \neq j'$, $\|\mathbf{U}_j^T \mathbf{U}_{j'}'\| = o_P(\frac{n}{m_n})$.

Démonstration Lemme 3.7 *Soit $j \notin S^*$.*

$$\begin{aligned} \|\mathbf{U}_j^T (\mathbf{Y} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n})\| &\leq \|\mathbf{U}_j^T (\mathbf{Y} - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{S_{\lambda_n}} - \mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r})\| \\ &\leq \|\mathbf{U}_j^T (\mathbf{Y} - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{S_{\lambda_n}})\| + \|\mathbf{U}_j^T \mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r}\| \\ &\leq \|\mathbf{U}_j^T (\mathbf{f}^*(\mathbf{X}) - \mathbf{U}_{S^*} \mathbf{b}_{S^*}^* + \underline{\epsilon} + \mathbf{U}_{S^*} \mathbf{b}_{S^*}^* - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{S_{\lambda_n}})\| + \|\mathbf{U}_j^T \mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r}\| \\ &\leq \|\mathbf{U}_j^T (\mathbf{f}^*(\mathbf{X}) - \mathbf{U}_{S^*} \mathbf{b}_{S^*}^*)\| + \|\mathbf{U}_j^T \Pi_j \underline{\epsilon}\| + \|\mathbf{U}_j^T (\mathbf{U}_{S^*} \mathbf{b}_{S^*}^* - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{S_{\lambda_n}})\| + \|\mathbf{U}_j^T \mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r}\|. \end{aligned}$$

Etudions chacun des 4 termes.

— *Etude de $\|\mathbf{U}_j^T (\mathbf{f}^*(\mathbf{X}) - \mathbf{U}_{S^*} \mathbf{b}_{S^*}^*)\|$:*

$$\begin{aligned} \|\mathbf{U}_j^T (\mathbf{f}^*(\mathbf{X}) - \mathbf{U}_{S^*} \mathbf{b}_{S^*}^*)\| &\leq \max_{j \notin S^*} \|\mathbf{U}_j^T \mathbf{U}_j\|^{1/2} \|(\mathbf{f}^*(\mathbf{X}) - \mathbf{U}_{S^*} \mathbf{b}_{S^*}^*)\| \\ &\leq O(1) \sqrt{\frac{n}{m_n}} m_n^{-\nu} \\ &\leq o(\lambda_n \sqrt{m_n}). \end{aligned}$$

— *Etude de $\|\mathbf{U}_j^T \Pi_j \underline{\epsilon}\|$:*

$$\begin{aligned} \|\mathbf{U}_j^T \Pi_j \underline{\epsilon}\| &\leq \max_{j \notin S^*} \|\mathbf{U}_j^T \mathbf{U}_j\|^{1/2} \|\Pi_j \underline{\epsilon}\| \\ &\leq \sqrt{\frac{n}{m_n}} o_P(\sqrt{m_n \log(nd)}) \\ &\leq o_P(\lambda_n \sqrt{m_n}). \end{aligned}$$

— *Etude de $\|\mathbf{U}_j^T (\mathbf{U}_{S^*} \mathbf{b}_{S^*}^* - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{S_{\lambda_n}})\|$:*

$$\begin{aligned} \|\mathbf{U}_j^T (\mathbf{U}_{S^*} \mathbf{b}_{S^*}^* - \mathbf{U}_{S^*} \hat{\mathbf{b}}_{S_{\lambda_n}})\| &\leq \max_{j \notin S^*} \|\mathbf{U}_j^T \mathbf{U}_j\|^{1/2} \|\mathbf{U}_{S^*}^T \mathbf{U}_{S^*}\|^{1/2} \|\mathbf{b}_{S^*}^* - \hat{\mathbf{b}}_{S_{\lambda_n}}\| \\ &\leq O(1) \sqrt{\frac{n}{m_n}} \sqrt{\frac{n}{m_n}} \frac{m_n}{\sqrt{n}} \\ &\leq O(1) \sqrt{n} \\ &\leq o(\lambda_n \sqrt{m_n}). \end{aligned}$$

— *Etude de $\|\mathbf{U}_j^T \mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r}\|$:*
 $\|\mathbf{U}_j^T \mathbf{U}_{j'}'\| = o_P(n/m_n),$

$$\begin{aligned}
\|\mathbf{U}_j^T \mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}\| &= \|\mathbf{U}_j^T \mathbf{U}_{S^*} \sum_{i=1}^{s^*} T_i^T T_i (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}\| \\
&\leq \sum_{i=1}^{s^*} \|\mathbf{U}_j^T \mathbf{U}_{S^*} T_i^T T_i (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}\| \\
&\leq \sum_{i=1}^{s^*} \|\mathbf{U}_j^T \mathbf{U}_{S^*} T_i^T\| \|T_i^T T_i\|^{1/2} \|(\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1}\| \|\underline{\mathbf{r}}\| \\
&\leq \sum_{i \in S^*} \|\mathbf{U}_j^T \mathbf{U}_i\| \|(\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1}\| \|\underline{\mathbf{r}}\| \\
&\leq o_P(n/m_n) \frac{m_n}{n} \lambda_n \sqrt{m_n} \\
&\leq o_P(\lambda_n \sqrt{m_n}).
\end{aligned}$$

Ceci permet de conclure que

$$\|\mathbf{U}_j^T (\mathbf{Y} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n})\| \leq o_P(\lambda_n \sqrt{m_n}).$$

□

Proposition 3.7

$$\|\mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}\| = O(1) \lambda_n \frac{m_n}{\sqrt{n}} = O(1) \sqrt{m_n \log(dm_n^2)}.$$

Démonstration Proposition 3.7

$$\begin{aligned}
\|\mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}\| &\leq \|\mathbf{U}_{S^*}^T \mathbf{U}_{S^*}\|^{1/2} \|(\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1}\| \|\underline{\mathbf{r}}\| \\
&\leq O(1) \lambda_n \sqrt{m_n} \sqrt{\frac{n}{m_n} \frac{m_n}{n}} \\
&\leq O(1) \lambda_n \frac{m_n}{\sqrt{n}} \\
&\leq O(1) \sqrt{m_n \log(dm_n^2)}.
\end{aligned}$$

□

Proposition 3.8

$$\tilde{\mathbf{b}}_{\lambda_n} - \underline{\mathbf{b}}_{S^*}^* = (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} (\mathbf{U}_{S^*}^T (\underline{\boldsymbol{\epsilon}} + B) - \underline{\mathbf{r}}),$$

et

$$\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \tilde{\mathbf{b}}_{\lambda_n} = (I - \Pi_{S^*})(\underline{\boldsymbol{\epsilon}} + B) + \mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \underline{\mathbf{r}}.$$

Démonstration Proposition 3.8 Nous étudions les deux expressions de la Proposition 3.8.

— Etude de $\tilde{\mathbf{b}}_{\lambda_n} - \underline{\mathbf{b}}_{S^*}^*$:

$$\begin{aligned}
\tilde{\mathbf{b}}_{\lambda_n} - \underline{\mathbf{b}}_{S^*}^* &= (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} (\mathbf{U}_{S^*}^T \underline{\mathbf{Y}} - \underline{\mathbf{r}}) - \underline{\mathbf{b}}_{S^*}^* \\
&= (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} (\mathbf{U}_{S^*}^T (\underline{\mathbf{Y}} - \mathbf{U}_{S^*} \underline{\mathbf{b}}_{S^*}^*) - \underline{\mathbf{r}}) \\
&= (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} (\mathbf{U}_{S^*}^T (\underline{\boldsymbol{\epsilon}} + B) - \underline{\mathbf{r}}).
\end{aligned}$$

— Etude de $\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \tilde{\mathbf{b}}_{\lambda_n}$:

$$\begin{aligned} \underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \tilde{\mathbf{b}}_{\lambda_n, \mathbf{S}^*} &= \underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \underline{\mathbf{b}}_{\mathbf{S}^*}^* + \mathbf{U}_{\mathbf{S}^*} \underline{\mathbf{b}}_{\mathbf{S}^*}^* - \mathbf{U}_{\mathbf{S}^*} \tilde{\mathbf{b}}_{\lambda_n} \\ &= \underline{\mathbf{B}} + \underline{\boldsymbol{\epsilon}} - \mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} (\mathbf{U}_{\mathbf{S}^*}^T (\underline{\boldsymbol{\epsilon}} + \underline{\mathbf{B}}) - \underline{\mathbf{r}}) \\ &= (\mathbf{I} - \Pi_{\mathbf{S}^*}) (\underline{\boldsymbol{\epsilon}} + \underline{\mathbf{B}}) + \mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}. \end{aligned}$$

□

Démonstration Lemme 3.8

$$\begin{aligned} \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\lambda_n}\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}_{\lambda_n}} \hat{\mathbf{b}}_{\mathbf{S}_{\lambda_n}}\|^2 &= \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \underline{\mathbf{b}}_{\lambda_n}\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}_{\lambda_n}} \underline{\mathbf{b}}_{\mathbf{S}_{\lambda_n}}\|^2 \\ &= \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} (\mathbf{U}_{\mathbf{S}^*}^T \underline{\mathbf{Y}} + \underline{\mathbf{r}})\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*}\|^2 \\ &= (\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} (\mathbf{U}_{\mathbf{S}^*}^T \underline{\mathbf{Y}} + \underline{\mathbf{r}}))^T (\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} (\mathbf{U}_{\mathbf{S}^*}^T \underline{\mathbf{Y}} + \underline{\mathbf{r}})) \\ &\quad - (\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*})^T (\underline{\mathbf{Y}} - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*}) \\ &= \|\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}\|^2 - 2(\underline{\mathbf{Y}} - \Pi_{\mathbf{S}^*} \underline{\mathbf{Y}})^T (\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}). \end{aligned}$$

□

Démonstration Lemme 3.9 D'après le Lemme 3.7

$$\|\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}\| = O(m_n \sqrt{n \log(m_n^2 d) / m_n} / \sqrt{n}) = O(\sqrt{m_n \log(m_n^2 d)}).$$

De plus,

$$\begin{aligned} (\underline{\mathbf{Y}} - \Pi_{\mathbf{S}^*} \underline{\mathbf{Y}})^T (\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}) &= (\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) + \underline{\boldsymbol{\epsilon}} - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*})^T (\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}) \\ &= (\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*})^T (\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}) + \underline{\boldsymbol{\epsilon}}^T (\Pi_{\mathbf{S}^*} \mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}) \\ &\text{car } \Pi_{\mathbf{S}^*} \mathbf{U}_{\mathbf{S}^*} = \mathbf{U}_{\mathbf{S}^*}. \end{aligned}$$

D'après Cauchy-Schwarz,

$$|(\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*})^T (\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}})| \leq \|\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*}\| \|\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}\|.$$

Nous avons

$$\|\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}\| = O(\sqrt{m_n \log(m_n^2 d)}),$$

et

$$\|\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*}\| \leq O_P(\sqrt{m_n}) \quad \text{d'après la Proposition 3.1,}$$

d'où

$$|(\underline{\mathbf{f}}^*(\underline{\mathbf{X}}) - \mathbf{U}_{\mathbf{S}^*} \hat{\mathbf{b}}_{\mathbf{S}^*})^T (\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}})| \leq O_P(m_n \sqrt{\log(m_n^2 d)}).$$

De même,

$$|\underline{\boldsymbol{\epsilon}}^T \Pi_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}})| \leq \|\Pi_{\mathbf{S}^*} \underline{\boldsymbol{\epsilon}}\| \|\mathbf{U}_{\mathbf{S}^*} (\mathbf{U}_{\mathbf{S}^*}^T \mathbf{U}_{\mathbf{S}^*})^{-1} \underline{\mathbf{r}}\|.$$

D'après le Lemme 3.3,

$$\|\Pi_{S^*} \underline{\epsilon}\| \leq o_P(\sqrt{m_n \log(nd)}),$$

d'où

$$|\underline{\epsilon}^T \Pi_{S^*} (\mathbf{U}_{S^*} (\mathbf{U}_{S^*}^T \mathbf{U}_{S^*})^{-1} \mathbf{r})| \leq o_P(m_n \sqrt{\log(nd) \log(m_n^2 d)}).$$

□

Démonstration Lemme 3.10

$$\begin{aligned} BIC(\lambda_n) - BIC(S^*) &= \log\left(\frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{\lambda_n} \tilde{\mathbf{b}}_{\lambda_n}\|^2}{\|\underline{\mathbf{Y}} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2}\right) \\ &= \log\left(1 + \frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{\lambda_n} \tilde{\mathbf{b}}_{\lambda_n}\|^2 - \|\underline{\mathbf{Y}} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2}{\|\underline{\mathbf{Y}} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2}\right) \\ &= \log\left(1 + \frac{\|\underline{\mathbf{Y}} - \mathbf{U}_{\lambda_n} \tilde{\mathbf{b}}_{\lambda_n}\|^2/n - \|\underline{\mathbf{Y}} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2/n}{\|\underline{\mathbf{Y}} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2/n}\right) \\ &= \log\left(1 + o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n})\right) \text{ car } \|\underline{\mathbf{Y}} - \mathbf{U}_{S_{\lambda_n}} \hat{\mathbf{b}}_{S_{\lambda_n}}\|^2/n \text{ borné d'après le Lemme 3.4} \\ &= o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n}) + o(o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n})). \end{aligned}$$

En effet,

$$m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n} \leq m_n \frac{\log(nd)}{n} = m_n \frac{\log(d)}{n} + m_n \frac{\log(n)}{n} \rightarrow 0.$$

Ensuite, nous montrons que $m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n}$ est négligeable par rapport à $\min(\log(2), v_n^2/m_n)$.

Comme $m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n} \leq \frac{m_n \log(nd)}{n}$, nous montrons plutôt que $\frac{m_n \log(nd)}{n}$ est négligeable par rapport à $\min(\log(2), v_n^2/m_n)$.

$\frac{m_n \log(nd)}{n} \rightarrow 0$, donc $\frac{m_n \log(nd)}{n} = o(\log(2))$, donc si $\min(\log(2), v_n^2/m_n) = \log(2)$, le résultat est immédiat. Montrons que $\frac{m_n \log(nd)}{n} = o(v_n^2/m_n)$.

D'après l'hypothèse 5.d,

$$\frac{m_n \log(nd)}{n} = o(\min_{j \in S^*} \|f_j\|^2),$$

et

$$\frac{v_n^2}{m_n} = \min_{j \in S^*} \|f_j\|^2 + o(\min_{j \in S^*} \|f_j\|^2),$$

donc

$$\frac{m_n \log(nd)}{n} = o\left(\frac{v_n^2}{m_n}\right),$$

et donc $m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n}$ est négligeable par rapport à $\min(\log(2), v_n^2/m_n)$.

□

Démonstration Théorème 3.4 Soit λ tel que $\text{card}(S_\lambda) \leq M$. Alors, par définition des MCO,

$$BIC(\lambda) \geq BIC(S_\lambda).$$

Supposons que $S^* \not\subseteq S_\lambda$, alors

$$\begin{aligned} BIC(\lambda) - BIC(\lambda_n) &\geq BIC(S_\lambda) - BIC(\lambda_n) \\ &\geq BIC(S_\lambda) - BIC(S^*) + o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n}) \quad \text{d'après le Lemme 3.9.} \end{aligned}$$

Or $BIC(S_\lambda) - BIC(S^*) \geq O(1)(\min(\log(2), \frac{v_n^2}{m_n}) + o(\frac{v_n^2}{m_n}))$,

et $o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n})$ négligeable par rapport à $\min(1, \frac{v_n^2}{m_n})$.

D'où, d'après le Théorème 3.1,

$$P(\min_{\lambda | \text{card}(S_\lambda) \leq M \text{ et } S^* \not\subseteq S_\lambda} BIC(\lambda) - BIC(\lambda_n) > 0) \xrightarrow{n \rightarrow +\infty} 1.$$

Supposons maintenant que $S^* \subseteq S_\lambda$,

alors

$$-BIC(\lambda) \leq -BIC(S_\lambda),$$

c'est-à-dire que

$$BIC(\lambda_n) - BIC(\lambda) \leq BIC(\lambda_n) - BIC(S_\lambda) \leq BIC(S^*) + o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n}) - BIC(S_\lambda).$$

Or $o_P(m_n \frac{\sqrt{\log(m_n^2 d) \log(nd)}}{n})$ négligeable par rapport à $\frac{m_n \log(nd)}{n}$.

D'où, d'après le Théorème 3.2,

$$P(\min_{\lambda | S^* \subseteq S_\lambda \text{ et } \text{card}(S_\lambda) \leq M} BIC(\lambda_n) - BIC(\lambda) < 0) \xrightarrow{n \rightarrow +\infty} 1.$$

Ceci démontre le théorème.

□



Perspectives

Dans cette conclusion, nous reprenons des perspectives d'étude déjà évoquées.

Perspectives méthodologiques

Comme le LASSO, le Group LASSO a tendance à ne sélectionner qu'un seul groupe de covariables lorsque ceux-ci sont fortement corrélés. Une pénalisation tenant compte de la corrélation peut être utilisée à place du Group LASSO. Par exemple, les auteurs de [43] combinent une pénalité L1 avec une pénalité tenant compte de la corrélation dans le contexte linéaire.

Nous utilisons l'algorithme de *Group Coordinate Descent* proposé par Brehenny *et al.* [19]. Qin *et al.* [97] et Scherrer *et al.* [101] proposent d'autres algorithmes pour résoudre le Group LASSO, dont certains peuvent être parallélisables.

L'ajout de nouvelles données pose la question de l'adaptativité des méthodes. Avec une bonne initialisation, les procédures peuvent être en partie adaptatives. De plus, il existe des algorithmes de descente de coordonnées qui s'apprennent en ligne (voir par exemple Wang et Banerjee, [117])

Nous avons travaillé sur les modèles additifs. Un travail similaire sur les modèles à coefficients variables peut être réalisé.

Perspectives applicatives

Etude nationale Nous avons travaillé avec des modèles par instant. Adapter nos procédures aux modèles à coefficients variables selon l'instant de la journée permet de s'affranchir de cette hypothèse et de travailler avec des modèles globaux, prenant ainsi en compte la corrélation entre les instants. Des hypothèses doivent alors être faites sur la matrice de covariance.

Une piste complémentaire d'amélioration est de rendre les méthodes adaptatives pour pouvoir les appliquer directement à la prévision court terme. Ceci pose deux questions :

- Les modèles additifs peuvent-ils facilement être adaptatifs ?
- Comment rendre les procédures de sélection adaptatives ?

Des études sont actuellement en cours à EDF pour répondre à la première question.

Enfin, les signaux de températures utilisés correspondent au signal de température nationale historique d'EDF, qui est une moyenne pondérée électriquement de 32 stations météorologiques. Au lieu de cette température nationale, nous pourrions utiliser les températures de chaque station météorologique directement en entrée des modèles ou les agréger, en utilisant par exemple la méthode proposée par Hong *et al.* [65].

Etude locale Les différents membres d'une maille d'un réseau interagissent. À court terme, utiliser l'information des charges consommées passées des autres zones ou postes sources peut améliorer les performances des modèles. Nous pouvons chercher à sélectionner les charges passées

des postes sources ou zones voisins avec les procédures Post2 pour améliorer la prévision court terme.

Les variables explicatives candidates sont très corrélées. Un changement de pénalisation, qui tiendrait compte de cette corrélation, pourrait améliorer les performances.

Pour ces données, nous avons travaillé sur la sélection des points de la maille du réseau où sont mesurées les variables de températures. Nous aurions pu créer à partir des températures mesurées à différents endroits un dictionnaire de covariables et faire le même type d'études que pour le portefeuille EDF.

Enfin, avec le développement de l'importance de la prise en compte des aléas, il est de plus en plus important d'avoir une vision probabiliste de la charge consommée. L'utilisation de la perte quantile [74] pénalisée pourraient être une piste à étudier.

Autres applications Dans l'article [70], Jollois *et al.* (2009) proposent d'utiliser les modèles additifs pour prévoir la pollution en Haute-Normandie. Ils utilisent un dictionnaire de covariables météorologiques comportant de la température, de l'humidité, de la pluie et du vent. Une étude similaire à l'étude du portefeuille EDF aurait pu être réalisée.

Dans une étude interne à EDF, des modèles additifs ont été appliqués à la prévision de production d'un champ d'éoliennes. L'objectif de l'étude est de trouver les signaux de vitesses du vent permettant la meilleure prévision de la production. Pour cela, le vent est mesuré le long d'une grille. Cette étude s'approche de l'étude réalisée sur la prévision de consommation locale.

Perspectives théoriques

Nous avons montré des propriétés de consistance pour un type d'estimateur en plusieurs étapes. L'étude de la procédure de type Post combinant le Group LASSO avec les P-Splines est plus compliquée, car l'estimateur dépend de la pénalisation associée à l'estimateur P-Splines.

Sous nos hypothèses, la méthode de démonstration utilisée pour le BIC ne permet pas de conclure des propriétés de consistance lorsque l'AIC est utilisé. En effet, sa pénalisation n'est pas suffisamment forte, ce qui nous empêche de conclure l'absence de faux négatifs lors de la première étape de la démonstration.

Une étude pour relâcher l'hypothèse de normalité des résidus pourrait être effectuée.

Annexes

Annexe : Vignette R CASA

Abstract

This document introduces the R package **CASA** which implements a statistical methodology for automatic selection and estimation of additive models. Specifically, the package includes functions which allow efficient selection and estimation in additive models, post-processing of selection and graphic representations. The methodology uses multi-stage estimators combining Group LASSO and P-Splines estimators. We illustrate its use with a real data example.

Keywords : Additive model, Multi-stage estimator, Penalized regression, R package, Variable selection

Introduction

Because they realize a good compromise between flexibility and complexity, additive models [57] are used in many fields of activity. For example, they have shown their efficiency in electricity load demand forecasting (see e.g. [94], [46], [55], [91], [39]) and in electricity prices forecasting (see e.g. [51]). Authors of [36] and [70] use them to study the air pollution. They are used by [14] for clinical prediction. In environmental context, they have been used to study the ground fish on the Northeast Coast of United State [93] and in the Bering sea [107] as well as on cod population dynamics [15]. To our knowledge, few studies and R packages devoted to component selection in additive models are available. We propose one method (see [10] and [111]) in the R package **CASA** which performs automatic component selection and estimation of additive models.

In Section A.1 we introduce the statistical framework and summarize corresponding R packages which rely on it. In Section A.2 we provide an overview on the R package **CASA** and develop its use on a real example.

A.1 Statistical framework and corresponding R packages

A.1.1 Additive models

We consider additive models as described in (A.1) :

$$Y_i = \beta_0 + f_1(X_{1,i}) + \cdots + f_d(X_{d,i}) + \varepsilon_i, \quad i = 1, \dots, n, \quad (\text{A.1})$$

where Y_i is an univariate response variable, $(X_{1,i}, \dots, X_{d,i})$ are the covariates influencing the response and ε_i is a random noise. The observed responses are supposed to be independent. The nonlinear functions $(f_1, \dots, f_d)$ are assumed to be smooth so that we can use a truncated series expansion into a B -splines basis to approximate them. With such an approximation, the model can be rewritten :

$$Y_i = \beta_0 + \sum_{j=1}^d \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(X_{i,j}) + \varepsilon_i, \quad i = 1, \dots, n,$$

where m_j is the dimension of the B -splines basis that models the effect f_j and $\{B_{j,k}^{q_j}(x), k = 1, \dots, m_j\}$ denote the corresponding normalized B -splines basis functions of order q_j . Bates and Venables' R package **splines** allows to compute B -splines.

Hastie's R package **gam** uses a backfitting algorithm to fit a generalized additive model (see [57]). More recently, in the R package **mgcv**, Wood [119] uses some penalized least squares method to fit generalized additive models and also some more general regression models.

Hereafter, we address the problem of selecting the nonzero additive component functions in such models and analyze the estimation of the resulting model (A.1). To achieve that, we use a penalized regression method which minimizes a criterion like :

$$Q^{OLS}(\boldsymbol{\beta}) + \sum_{j=1}^d p_{\lambda_j}(\boldsymbol{\beta}_j),$$

where

$$Q^{OLS}(\boldsymbol{\beta}) = \sum_{i=1}^n \left(Y_i - \beta_0 - \sum_{j=1}^d C_{ij}(\boldsymbol{\beta}_j) \right)^2,$$

with $C_{ij}(\boldsymbol{\beta}_j) = \sum_{k=1}^{m_j} \beta_{j,k} B_{j,k}^{q_j}(X_{i,j})$ and p_{λ} some appropriate penalty function with a penalty parameter λ .

In the generalized linear context, Hastie and Efron propose, in the R package **lars**, some stepwise, LASSO and least angle regression methods (see [40]). In the same context, Friedman *et al.* propose, in the R package **glmnet**, a covariate selection and a model estimation based on LASSO and the elastic net (see [124]). In **lqa**, Ulbricht [114] proposes to fit GLMs based on penalized maximum likelihood inference.

We consider nonlinear additive models. As each covariate effect is modeled using a B -splines projection, coefficients of our additive models are naturally grouped, each group corresponding to one covariate. The common tool to deal with that a situation is the Group LASSO estimator introduced by [122]. Meier proposes in the R package **grplasso** a Group LASSO algorithm and applies it in the particular context of logistic models (see [88]). In the R package **grpreg**, Brehenny proposes several penalized based efficient algorithms which include group selection methods such as Group LASSO, Group MCP, and Group SCAD as well as bi-level selection methods such as the group exponential LASSO, the composite MCP, and the Group Bridge (see [66] for more details). In **mgcv**, Wood proposes some appropriate modifications in his adopted P -splines penalty to perform component selection (see [86]). However, we have seen on simulations that this method can have many false positives and that zero effect are never strictly thresholded to zero.

A.1.2 Multi-stage estimator : the Post2 procedure

As LASSO, Group LASSO is a powerful selection procedure similar to soft thresholding in the orthogonal design case but incurs poor performance in prediction because of the bias induced by such soft thresholding. In contrast, P -splines estimators (see for example the Ramsey and Ripley's R package **pspline** in univariate context) realize good forecasting performances as smoothing prevents from overfitting but P -splines do not achieve any feature selection. Belloni *et al.* [16] combine LASSO with OLS estimator in linear context and obtain good performances on synthetic data. Inspired from their work, we propose in [10] several procedures and retain the one called the Post2 procedure that leads to good forecasting performances by combining the Group LASSO and the P -splines. We give below a detailed description of this algorithm. We denote Λ_{GrpL} a uniform grid between $\lambda_{\min} = 0$ and $\lambda_{\max}$ the value of λ for which all the effects are shrunk to 0. In the following, we will call Post2Bic, Post2Aic and Post2Gcv the Post2 procedure associated to BIC [102], AIC [1] or GCV [115] criteria respectively.

Algorithm Post2**1. First step : Subset selection (Group LASSO)**For each $\lambda_i \in \Lambda_{GrpL}$

— Solve

$$\hat{\beta}^{\lambda_i} = \arg \min \{ Q^{OLS}(\beta) + \lambda_i \sum_{j=1}^d \sqrt{m_j} \|\beta_j\|_2 \}$$

— Denote $S^{\lambda_i} = \{j \mid \hat{\beta}_j^{\lambda_i} \neq 0\}$ **2. Second step : Estimation of additive models (P-splines)**For each support set $S^{\lambda_s} \in \{S^{\lambda_{min}}, \dots, S^{\lambda_{max}}\}$

— Compute

$$Q_{S^{\lambda_s}}^{OLS}(\beta) = \sum_{i=1}^n \left(Y_i - \beta_0 - \sum_{j \in S^{\lambda_s}} C_{ij}(\beta_j) \right)^2$$

— Solve

$$\tilde{\beta}^{S^{\lambda_s}} = \arg \min \{ Q_{S^{\lambda_s}}^{OLS}(\beta) + \sum_{j \in S^{\lambda_s}} \mu_j \|D_2 \beta_j\|_2^2 \},$$

where the penalty parameters $(\mu_j)_{j \in S^{\lambda_s}}$ are selected by numerical minimization of GCV. D_2 denotes the finite difference operator of order 2.— Compute a model selection criterion (MSC), like BIC, AIC or GCV for each $\tilde{\beta}^{S^{\lambda_s}}$ **3. Third step : Selection of the final model** Select $\tilde{\beta}^{S^{\lambda_b}}$ which minimizes the chosen MSC**A.2 Overview of the R package CASA**

The R package **CASA** (Component Automatic Selection in Additive models) implements the Post2 algorithm described above. It relies on the packages listed in Table A.1.

Package	Description
mgcv	GAM estimation with penalized methods
grpreg	Group LASSO estimator computation
splines	Splines computation
magic	matrix diagonalization
abind	multi-dimensional vectors combination

TABLE A.1 – R packages to import

We quickly describe some functions of the package which are used on the next example :

— **Estimation :**

— **by using Algorithm Post2 :** *Post2*, *Post2TestingSet* or *Post2Multiple*, the second is used when there is a validation/testing set and the third in case of multiple series.

— **by using a *Post2* object :** *UsePost2Selection*. Use it when some extensions offered by *gam* of **mgcv** are desired.

— **by using a list of *Post2* or *Post2TestingSet* objects :** *postProcessing*, which allows a post-processing of selected covariates subsets. The function uses *cleaningSelection*, which takes a list of *Post2* or *Post2TestingSet* objects or a *Post2Multiple* object and computes the post-proceeding candidate covariates subsets.

— **Prediction a *Post2* object :** *Post2Predict*

— **Graphics :**

— **Selection :** *plotSelection* takes a *cleaningSelection* object, and plots what covariates

are selected for each instant (can be only one instant too, and in this case, that gives the covariates selected by a *Post2* object) and the frequency of selection for each covariates along the choosen granularity.

- **Forecasting performances** : *plotForecasting* takes a list of *Post2TestingSet* objects and of some benchmarks forecasting, and plots the sorted MAPE and RMSE for each of them and the MAPE and the RMSE in function of some choosen categorical covariates.

To further proceed, we start, in subsection A.2.1, by presenting a case study that allows us to show some of the possibilities of **CASA**. Then, in subsection A.2.2, we describe more precisely the function *Post2* and display the activity diagram of the function *mainPost2*, since it summarizes some of the possibilities offered by **CASA**. In subsection A.2.3 we address some questions and the relevant answers about the **CASA** use. Subsection A.2.4, is devoted to the presentation of others functions and algorithms implemented in the package and addresses some perspectives.

A.2.1 A real data example of **CASA**'s use

A.2.1.1 The data

The Global Energy Forecasting Competition presented in [63] aims at forecasting and backcasting hourly loads (in kW) of 20 US zones corresponding to points of the US grid. The participants were provided with past load of the 20 zones and the recordings of temperature of 11 stations data between January 2004 and June 2008. No information was provided about the geographical localization neither the grid connections and this lack of information induced one of the major issues that participants were faced with :

- How many meteorological stations to use ?
- What stations for each zone ?

The fourth first teams of GEFCom 2012 assume a constant number of meteorological stations (five for Charlton and Singleton [105], eleven for Lloyd [84] and one for Nedellec *et al.* [91] and Ben Taieb and Hyndman [17]). As suggested by [65], one should improve the forecasting performance by relaxing this hypothesis. Actually, as illustrated by Figure A.1, each zone could potentially cover an unknown and potentially large territory due to the grid topology and participants had to develop modeling strategies in order to recover this hidden structure. We think that this public dataset and in particular this problem of stations selection makes a nice case study for illustrating our method.

We divide the dataset into a training set : 2004-2007 and a testing set : 2008. Figure A.2 represents the load curves of the different zones. It is quite clear that zones 9 and 10 have very different patterns as already pointed by the participants. We exclude it from our analysis as well as zone 7, because it is a translation from the zone 3.

We are presenting here the use of the R package **CASA** on zone 1. We start by the study of mid-day data.

A.2.1.2 Modelling a single time serie : zone 1 at mid-day

In this subsection we forecast the load demand at mid-day for the zone 1. We present 3 functions. *Post2* implements the Algorithm Post2. The user can choose the model selection criterion (BIC, AIC and GCV). *gam* function of package **mgcv** allows several extensions which are not addressed in *Post2*. The function *UsePost2Selection* allows to use a covariates subset selected by *Post2* in a *gam* object. It takes the selected covariates by *Post2* and integrates them in the *gam* formula chosen by the user, thus retrieving all the possibilities offered by *gam* of **mgcv**. The function *Post2TestingSet* has to be used whenever there is a testing (or validation) set available. It returns forecasting results and forecasting performance criteria in addition to the objects returned by *Post2*.

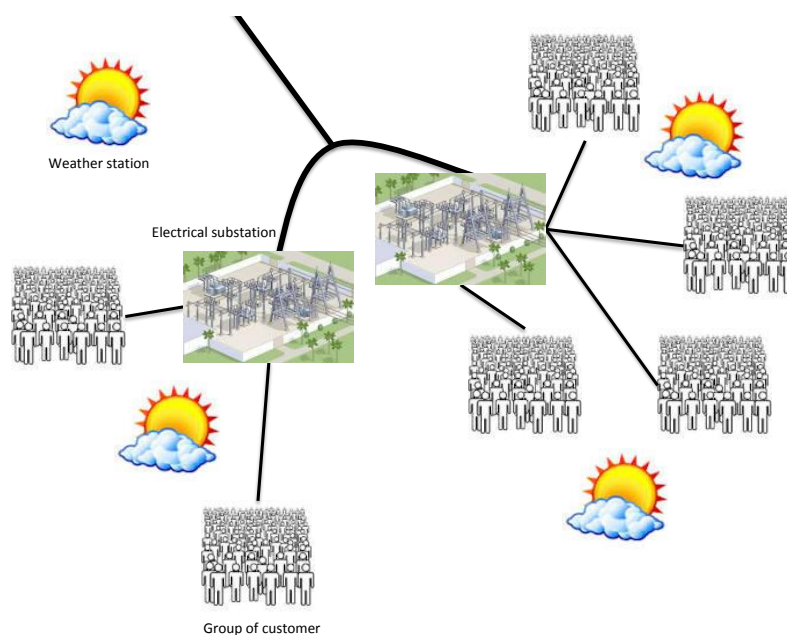


FIGURE A.1 – GEFCOM 2012 illustration

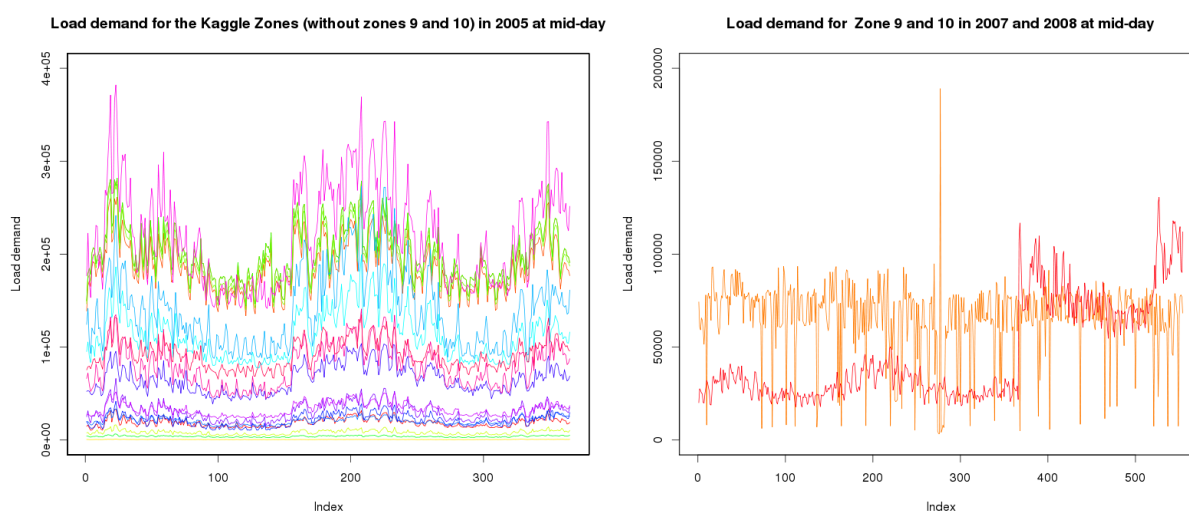


FIGURE A.2 – Load demand of GEFCOM 2012

After the above description, let us build the training and testing sets :

```
> library("casa")
> load("VignetteData.RData")
> dataTrain=VignetteData$Data0
> dataTest=VignetteData$Data1                                #Testing set
> d2=dataTrain[which(dataTrain$Hour==12),]                    #Mid-day
> Y=dataTrain[which(dataTrain$Hour==12),]$Zone1               #Response covariate
> training=cbind(d2$Station1, d2$Theta1
+ , d2$Station2, d2$Theta2, d2$Station3, d2$Theta3
+ , d2$Station4, d2$Theta4, d2$Station5, d2$Theta5
+ , d2$Station6, d2$Theta6, d2$Station7, d2$Theta7
+ , d2$Station8, d2$Theta8, d2$Station9, d2$Theta9
+ , d2$Station10, d2$Theta10, d2$Station11, d2$Theta11
+ )                                                            #Candidates covariates
>
```

We assume the full model described by (A.2) :

$$Y_t = \sum_{u=1}^{11} f_{1,u}(\theta_{t,u}) + \sum_{u=1}^{11} f_{2,u}(T_{t,u}) + \varepsilon_t, \quad (\text{A.2})$$

where

- Y_t is the response variable at time t , the load in our case. It corresponds to Y in the above R-code.
- $T_{t,u}$ and $\theta_{t,u}$ are the temperature and the smoothed temperature of the meteorological station u at time t . It corresponds to $Station_u$ and $Theta_u$ respectively.

We aim at selecting models, without considering the geographical aspect, that means $T_{t,2}$ can be selected and not $\theta_{t,2}$ for example. We may apply the function *Post2* :

```
> model1=Post2( Y=Y, training=training,group=c(1:22),BS="ps", k=10, criterion=0)
```

where *group* runs from 1 to 22 because there are 22 covariates, *BS* corresponds to the splines basis which are used for the estimation ("ps" is *P*-splines), *k* is the degree of freedom and *criterion* corresponds to the model selection criteria used. If zero, then BIC, GCV and AIC are used. *Post2* returns a list of objects with as many *gam* objects and selected covariates subsets as there are of model selection criteria :

```
> model1$Post2Bic
```

Family: gaussian

Link function: identity

Formula:

```
X1 ~ s(X2, bs = "ps", k = 10) + s(X3, bs = "ps", k = 10) + s(X4,
    bs = "ps", k = 10)
```

```
<environment: 0x7eff93de5db0>
```

Estimated degrees of freedom:

```
8.53 5.45 3.63 total = 18.6
```

GCV score: 3581964

```
> model1$VarPost2Bic
```

```
[1] 4 11 19
```

Similar outputs hold for `Post2Gcv` and `Post2Aic`. For `Post2Bic` the covariates 4, 11 and 19 have been selected which means that the smoothed temperature of station 2 and the real temperatures of stations 6 and 10 have been selected. We can now use the function `Post2Predict` to forecast a testing set (build this subset by substituting `d2` with `d3` in the previous code) :

```
> Forecasting1=Post2Predict(post2Object=model1, testing=testing)
```

The function returns as many forecast as there are models in the `Post2` object. As said at the beginning, one important issue of GEFCom 2012 was to choose a number of meteorological stations and to select them. We would like that the procedures select or not simultaneously the real and the smoothed temperatures of each station. To do this we use the parameter `group` of the function `Post2` as shown next :

```
> group1=rep(1:11, each=2)
> model2=Post2(Y=Y, training=training, group=group1, BS="'ps'")
```

`group1` has the given form because we want to organise the covariates two by two. In this case `VarPost2Bic`, `VarPost2Gcv` and `VarPost2Aic` return the covariates groups selected by `Post2Bic`, `Post2Gcv` and `Post2Aic` respectively. This means that if `VarPost2Bic` returns 10 for example, then both real and smoothed temperatures of station 10 have been selected.

We are now interested in the selection of weather stations. However the load demand is generally dependent of calendar covariates. We would like to introduce calendar covariates and we need to introduce continuous (e.g. time of the year) or factorial covariates (e.g. day type). We work on the new saturated model (A.3) :

$$Y_t = \sum_{q=1}^7 m_q \mathbb{1}_{DayType_t=q} + \sum_{u=1}^{11} f_{1,u}(\theta_{t,k_u}) + \sum_{u=1}^{11} f_{2,u}(T_{t,k_u}) + f_3(Toy_t) + f_4(t) + \varepsilon_t, \quad (\text{A.3})$$

where `DayType` is the day type and `Toy` goes from 0 to 1 from the beginning to the end of year. We do address the calendar covariates selection since we assume there are significant in all models.

```
> trainingNum=as.matrix(cbind(d2$Time,d2$Toy))
> trainingFact=cbind(d2$DayType)
> model3=Post2( Y=Y,training=training
+               ,group=group1,BS="'ps'", k=10
+               ,trainingNum=trainingNum,trainingFact=trainingFact
+               ,kVCNP=c(4,30),BS2=c("'ps'", "'ps'"),criterion=0)
> testingNum=as.matrix(cbind(d3$Time,d3$Toy))
> testingFact=cbind(d3$DayType)
> Forecasting2=Post2Predict(post2Object=model3, testing=testing
+                           , testingNum = testingNum, testingFact = testingFact)
```

`kVCNP` and `BS2` are the degree of freedom and the type of splines basis which are associated with `toy` and `t`, which are continuous. $kVCNP \geq 4$ and `BS2` is a vector of `"ps"` or `"cr"`. In function

Post2, the user can specify *model* = 0. Then the models are not fitted and the function returns only the selected covariates.

All the possibilities offered by *gam* of **mgcv** are not used by *Post2* so we propose a function *UsePost2Selection* to allow the use of *Post2* selected covariates to integrate them in the function *gam*. For example, to add a linear trend in our model we can use a formula syntax described below :

```
#Covariates in all selected subsets
> trainingKernelModel=cbind(d2$Time, d2$Toy, d2$DayType)
> #name of previous covariates
> kernelModelCovName=c( "Time", "Toy", "DayType")
> #New formula
> kernelModelFormula="Time+s(Toy,k=30, bs='cc')+as.factor(DayType)"
> #Splines basis of candidates covariates
> BS3=c(rep("'ps'", 22))
> #degree of freedoms of splines basis of candidates covariates
> k3=rep(c(10,5), 11)
> model4=UsePost2Selection(post2Object=model3,Y=Y,training=training
+                           ,trainingKernelModel=trainingKernelModel
+                           ,kernelModelFormula=kernelModelFormula
+                           ,kernelModelCovName=kernelModelCovName
+                           ,BS3=BS3,k3=k3
+                           )
> model4$Post2Bic
```

Family: gaussian

Link function: identity

Formula:

```
Load ~ Time + s(Toy, k = 30, bs = "cc") + as.factor(DayType) +
      s(X3, bs = "ps", k = 10) + s(X4, bs = "ps", k = 5) + s(X19,
      bs = "ps", k = 10) + s(X20, bs = "ps", k = 5)
<environment: 0x7eff93b9d120>
```

Estimated degrees of freedom:

```
18.97  5.53  3.54  5.50  3.34  total = 44.88
```

fREML score: 12582.36

Assume now we study in a case where we have at our disposal a training and a testing or validation set. Then we can use the function *Post2TestingSet*, which allows, in addition to the models fitted and the covariates selected, to return the forecasting MAPE and RMSE. *Post2TestingSet* can be used as :

```
> kVCNP=c(4,30)
> BS2=c("'ps'", "'ps'")
> BS="'ps'"
> k = 10
> model5=Post2TestingSet(Y=Y, Yt=Yt, training=training
+                       , testing=testing, group=group1, BS=BS, k = 10
+                       , trainingNum = trainingNum
```

```

+           , trainingFact = trainingFact
+           , testingNum = testingNum
+           , testingFact = testingFact
+           , kVCNP = kVCNP, BS2 = BS2
+         )

```

In the following we will study zone's 1 load demand for all day, instead of mid-day.

A.2.1.3 Modelling multiple time series : zone's 1 all day load demand

We develop here our approach on a multiple time series composed of 24 electricity load curves (one time series per hour of the day). Decomposing the electricity load in this way is a standard approach (see e.g. [46]; [94]) and is used because the covariates effects and the load are usually strongly correlated with the hour. We can easily imagine other settings where such a multiple approach could be used. For example, in a telecommunication context, one could imagine similar type of assumptions for forecasting the Internet flow or the number of text messages exchanged; or in a transport context, for forecasting the number of passagers of a bus, a taxi trajectory or the time of travel. The function *Post2Multiple* returns a list of *Post2* or *Post2TestingSet* objects for each instant.

```

> training2=cbind(dataTrain$Station1, dataTrain$Theta1, dataTrain$Station2
+               , dataTrain$Theta2,dataTrain$Station3, dataTrain$Theta3
+               , dataTrain$Station4, dataTrain$Theta4,dataTrain$Station5
+               , dataTrain$Theta5, dataTrain$Station6, dataTrain$Theta6
+               , dataTrain$Station7, dataTrain$Theta7,dataTrain$Station8
+               , dataTrain$Theta8,dataTrain$Station9, dataTrain$Theta9
+               , dataTrain$Station10,dataTrain$Theta10
+               , dataTrain$Station11, dataTrain$Theta11)
> trainingNum2=as.matrix(cbind(dataTrain$Time,dataTrain$Toy))
> trainingFact2=cbind(dataTrain$DayType)
> testing2=cbind(dataTest$Station1, dataTest$Theta1, dataTest$Station2
+               , dataTest$Theta2,dataTest$Station3, dataTest$Theta3
+               , dataTest$Station4, dataTest$Theta4,dataTest$Station5
+               , dataTest$Theta5, dataTest$Station6, dataTest$Theta6
+               , dataTest$Station7, dataTest$Theta7,dataTest$Station8
+               , dataTest$Theta8,dataTest$Station9, dataTest$Theta9
+               , dataTest$Station10, dataTest$Theta10
+               , dataTest$Station11, dataTest$Theta11)
> testingNum2=as.matrix(cbind(dataTest$Time,dataTest$Toy))
> testingFact2=cbind(dataTest$DayType)
> seriesDividingCov=dataTrain$Hour
> seriesDividingCovTesting=dataTest$Hour
>
> model6=Post2Multiple(seriesDividingCov=seriesDividingCov, Y=dataTrain$Zone1
+               , training=training2, group=rep(1:11, each=2)
+               , BS="'ps'", trainingNum = trainingNum2
+               , trainingFact = trainingFact2
+               , seriesDividingCovTesting=seriesDividingCovTesting
+               , Yt=dataTest$Zone1,testing=testing2
+               , testingNum = testingNum2, testingFact = testingFact2
+               , kVCNP = kVCNP, BS2 = BS2)

```

The selected covariates subsets change at each instant. We focus on the representation of these subsets (see the function *plotSelection*). Moreover, as different instants are strongly correlated, we implement a post-processing algorithm to improve the selection and to select only common selected subsets (see the functions *cleaningSelection* and *PostProcessing*). We implement *plotForecasting* function which takes as inputs a list of *Post2TestingSet* objects and some benchmarks and plots different figures estimating the forecasting performances.

First, we are interested in selecting covariates subsets. The questions are :

- What are the covariates selected for each instant ?
- What is the selection frequency of each covariate ?

The function *cleaningSelection* takes as inputs a list of *Post2* or *Post2TestingSet* objects or a *Post2Multiple* object and returns a boolean matrix whose entries indicate for what instants a covariate is selected or not. It takes as inputs some vectors which correspond to the percentage of minimum presence of covariates for each method and returns a matrix indicating what covariates satisfy this minimum percentage of selection. The function *plotSelection* allows to plot a *cleaningSelection* object. For legibility, we advise to use *criterion* = 4, 5 or 6 for this function. Let us give an example of their use :

```
###BICGrid = c(5,50,75) means that for Post2Bic,
### covariates selected for at least
### 5, 50 and 75 instant percent respectively are shown.
> cleaning1=cleaningSelection(listPost2=model6
+                               , BICGrid = c(5,50,75)
+                               , AICGrid = c(5,10,20,90)
+                               , GCVGrid = c(1,5,15,50,75)
+                               )
> plotSelection(cleaning1,groupCov=1:11
+               , criterion = 4
+               , yaxisName = 1:11
+               , plot.SelectedCov = 1
+               , plot.PostProcessing1 = 1
+               , xlab="Hour"
+               , selectedCovTitle="Covariates selected by Post2Bic"
+               , postProcessing1Title="Covariates along the percentage"
+               , xaxisBicName=c(5,50,75)
+ )
```

The resulting figures are given next. On these figures, “Hour” is the studied time granularity. “Covariate number” corresponds to the label used for each covariate. “Percentage” corresponds to the selection minimum percentage. There is a point when a covariate is selected for an instant or for a percentage. For example, station 10 is selected at 12h but not at 1h and is selected for at least 50% of hours but not for at least 75% of hours for Post2Bic. The user can choose to not plot all the covariates by changing *groupCov* parameter.

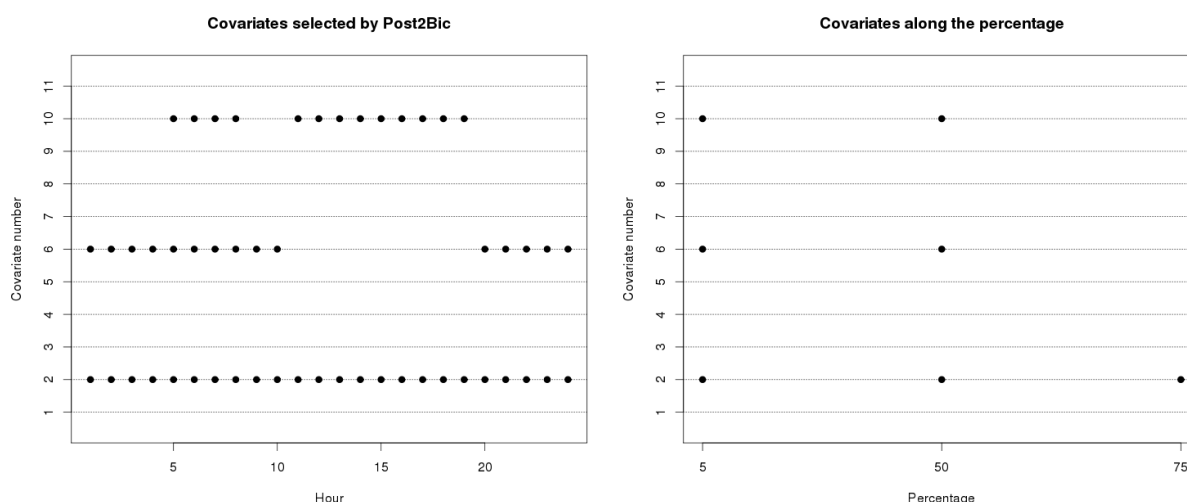


FIGURE A.3 – Instant where stations are selected FIGURE A.4 – Covariates selection frequency

Candidate variables subsets can be obtained by considering the subsets of covariates given in Figure A.4. Function *postProcessing* allows to do some post-processing treatment. The user can choose two methods. In the first, the function returns the estimating MAPE and RMSE of each candidate and the user can apply an “elbow” method [110]. In the second, which is automatic, the user divides the training set in two subsets, with a first subset where the models are training and a second subset where the final model is selected by minimizing the MAPE or the RMSE.

```
> Y=dataTrain$Zone1
> trainingKernelModel=as.matrix(cbind(dataTrain$Time,dataTrain$Toy,dataTrain$DayType))
> model7=postProcessing(listPost2=model6,BICGrid = c(5,60,75)
+                        , criterion = 4, Y=Y
+                        , training=training2
+                        , trainingKernelModel=trainingKernelModel
+                        , kernelModelFormula=kernelModelFormula
+                        , kernelModelCovName=kernelModelCovName
+                        , BS3=BS3, k3=k3
+                        , seriesDividingCov=seriesDividingCov
+                        , trainingValidation = 1
+                        , trainingVector = c(which(dataTrain$Year%in%c(2004:2006)))
+                        , validationVector =c(which(dataTrain$Year==2007))
+                        , trainingValidationCriterion = 1)
> model7$ModeleP2BicCorrige[[1]]
> model7$ModeleP2BicCorrige[[2]]

> model7$ModeleP2BicCorrige[[1]]
Family: gaussian
Link function: identity

Formula:
Load ~ Time + s(Toy, k = 30, bs = "cc") + as.factor(DayType) +
      s(X3, bs = "ps", k = 10) + s(X4, bs = "ps", k = 5) + s(X11,
      bs = "ps", k = 10) + s(X12, bs = "ps", k = 5) + s(X19, bs = "ps",
      k = 10) + s(X20, bs = "ps", k = 5)
<environment: 0xd827a40>
```

Estimated degrees of freedom:

```
16.37  6.01  3.83  1.00  2.88  5.20  3.30
total = 46.58
```

fREML score: 12088.69

```
> model7$ModeleP2BicCorrige[[2]]
```

Family: gaussian

Link function: identity

Formula:

```
Load ~ Time + s(Toy, k = 30, bs = "cc") + as.factor(DayType) +
  s(X3, bs = "ps", k = 10) + s(X4, bs = "ps", k = 5) + s(X11,
  bs = "ps", k = 10) + s(X12, bs = "ps", k = 5) + s(X19, bs = "ps",
  k = 10) + s(X20, bs = "ps", k = 5)
<environment: 0xd827a40>
```

Estimated degrees of freedom:

```
15.67  5.11  3.81  5.52  1.42  5.17  3.61
total = 48.32
```

fREML score: 12085

With a list of *Post2TestingSet*, we get a lot of information that the function *plotForecasting* allows to retrieve. In addition to the figures obtained with *plotSelection* function, this function plots the forecasting MAPE and RMSE of each model and benchmark in an ascending order. It plots the forecasting MAPE and RMSE of each model and benchmark as a function of instants and more generally as a function of each factorial covariate desired (e.g. the month) too. On our example, this gives :

```
> benchmarkMatrix = cbind(dataA[which(dataA$Year==2007)[1:dim(dataTest)[1]],]$Zone1
>      ,dataA[which(dataA$Year==2006)[1:dim(dataTest)[1]],]$Zone1)
> plotForecasting(listPost2TestingSet=model6,Yt=dataTest$Zone1
>      , seriesDividingCovTesting=seriesDividingCovTesting
>      , groupCov=1:11
>      , criterion = 4, yaxisName = 1:11
>      , titleMape = "MAPE", titleRmse = "RMSE"
>      , benchmarkMatrix = benchmarkMatrix
>      , benchmarkName = c("Past Load 1", "Past Load 2")
>      , colForecasting = c("red", "green4", "blue")
>      , lwdForecasting = 5
>      , mapeDividingTitle = "MAPE by Instant"
>      , rmseDividingTitle = "RMSE by Instant"
>      , factCov = list(dataTest$Month)
>      , mapeFactCovTitle = "MAPE by Month"
>      , factCovTitle = list()
>      , xaxisFactCov = "Month"
>      , rmseFactCovTitle = "RMSE by Month")
```

and we obtain, in addition to Figures A.3 and A.4, the next Figure :

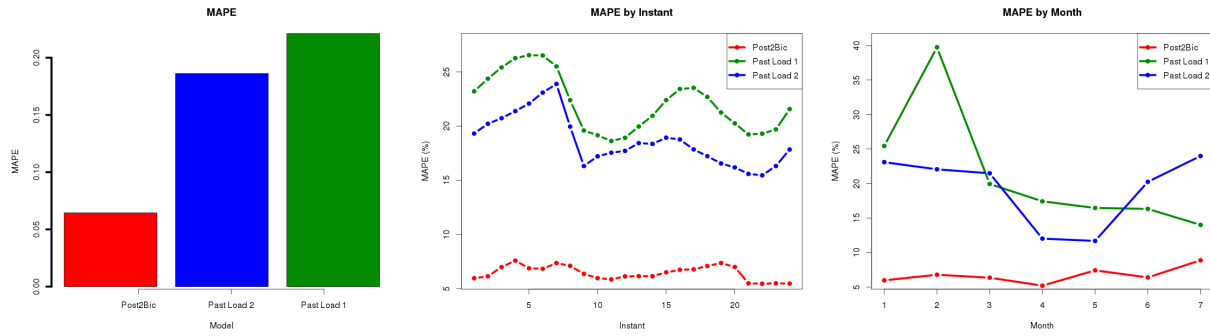


FIGURE A.5 – Forecasting performance

A.2.2 Two main functions

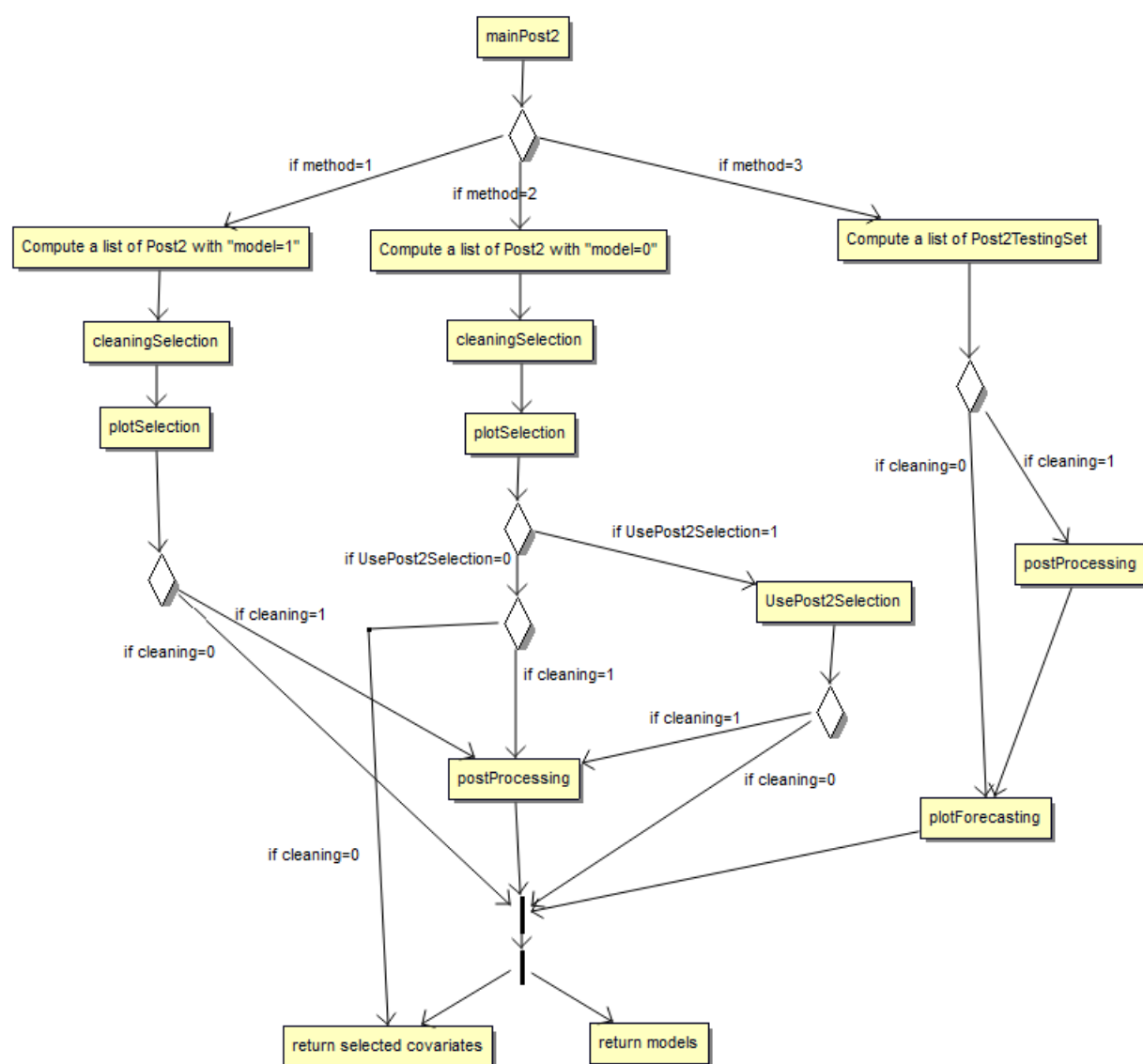
In this subsection we explain the inputs of *Post2* and the *mainPost2* activity diagram.

A.2.2.1 *Post2*

The function *Post2* uses the functions *grpreg* and *gam* of **grpreg** and **mgcv** packages. Next, we describe the inputs of the function. *Y* corresponds to the response variable. *training* corresponds to the matrix of candidate covariates. *group* allows to group the selection of some covariates. *BS* is the common splines used to model the effect of candidate covariates in the second step of Algorithm *Post2*. We may use *P*-splines or cubic regression splines (see [119]) when "*ps*" or "*cr*" are specified respectively. Caution : the user must necessarily specify " ' Splines basis ' ". *k* is the degree of freedom of the splines used to approximate the effect of candidate covariates. It is the same for all the splines, because, as the Group LASSO penalty and the model selection criterion (BIC, GCV or AIC) are weighted by the dimension of each group, it will be unfair to use different degrees of freedom. Sometime, we may want to integrated some extra continuous or factorial covariates. They can be included in *trainingNum* and in *trainingFact*. The inputs *max.iter*, *eps* and *nlambda* concern some parameters of *grpreg*. The input *model* is a dummy covariate which indicates if final models need to be fitted or not. User specifies in *criterion* the model selection criterion which will be used. The function returns the final model(s) and the selected subsets.

A.2.2.2 *mainPost2*

The function *mainPost2* allows to automatize the sequence of several functions. Figure A.6 shows its activity diagram.

FIGURE A.6 – Activity diagram of *mainPost2*

A.2.3 Questions and Answers

Assume we study an additive model...

- **How can we select covariates?** Use the functions *Post2* or *Post2TestingSet* (if there is a validation or testing set), *Post2TestingSet* or *Post2Multiple* (in multiple series context)
- **How can we simultaneous select two covariates?** Allocate them the same number in the parameter *group* of the functions *Post2*, *Post2TestingSet* or *Post2Multiple*
- **How can we add some covariates which are compulsorily selected?** Depending of the type of covariates (continuous or categorical), add it in the appropriate parameter in the functions
- **How can we use more possibilities offered by the function *gam* of *mgcv*?** Use the function *UsePost2Selection*
- **Assume we are working on a study where the experiment can cut along a categorical variable. How can we homogenize the selected covariates subsets?** Use *cleaningSelection* and *PostProcessing*
- **How can we represent the selected covariates subsets?** Use *plotSelection*
- **How can we represent the forecasting performances?** Use *plotForecasting*
- **How can we automatize the study of the additive model?** Use *mainPost2*

A.2.4 Perspectives

We have not described all the possibilities and functions available in the package. One may apply the selection algorithm in a linear context by combining OLS and (Group) LASSO estimators. Another function will be soon available implementing the algorithm called Ante2 in [10], which consists in selecting the final subset before the *P*-splines step. Important upgrades of the package will be the parallelization of some functions, the use of an other penalization procedures different than Group LASSO (to take into account a possible strong correlation between covariates), the use of more possibilities available in **mgcv** and the development of some adaptive algorithms.

A.3 *Post2* help

Hereafter is the help manual of the function *Post2*.

R documentation

of 'Post2.Rd'

June 8, 2015

Post2	Post2 Procedure
-------	-----------------

Description

Estimate a multi-step estimator combining Group LASSO and P-Splines

Usage

```
Post2(Y, training, group, BS, k = 10
, trainingNum = NULL, trainingFact = NULL
, kVCNP = NULL, BS2 = NULL, dfCalculation = 1
, max.iter = 20000, eps = 0.005, nlambda = 100
, model = 1, criterion = 0)
```

Arguments

Y	Vector, response variable
training	Matrix or dataframe, design with covariates we search to select
group	Vector, consecutive integers describing the grouping of the coefficients. Integers for penalized groups must run from 1 to the number of penalized groups.
BS	"ps" or "cr", Spline associated to training
k	Integer, >4, number of degree of freedoms for the spline associated to training, typically k=10
trainingNum	Matrix or dataframe, Design with the continuous covariates which are in all models. Default is NULL
trainingFact	Matrix or dataframe, Design with the indicator covariates. Default is NULL
kVCNP	Vector of integer >3 of size dim(trainingNum)[2], degree of freedom associated at each covariates of trainingNum
BS2	Vector of "ps" or "cr", Spline associated to each covariates of trainingNum

dfCalculation

1 or 2, method use for computing the degree of freedom of the models, 1 counts the number of non zero, 2 is less strict. Default is 1

max.iter integer, maximum of iteration for the Group LASSO. Default is 20000

eps Real, Convergence threshold. The algorithm iterates until the relative change in any coefficient is less than eps. Default is 0.005

nlambda Integer. The number of lambda values for Group LASSO. Default is 100

model 0 or 1. If 1, then the function computes the models, if 0, the function returns only the covariates selected. Default is 1

criterion Numeric. If ==0 if BIC, AIC and GCV are used, 1 if BIC and AIC are used, 2 if BIC and GCV are used, 3 if AIC and GCV are used, 4 if BIC is used, 5 if AIC is used, 6 if GCV is used

Value

Post2Bic	Gam object with the model obtaining with Post2Bic
Post2Aic	Gam object with the model obtaining with Post2Aic
Post2Gcv	Gam object with the model obtaining with Post2Gcv
VarPost2Bic	Selected covariate (by group) given by Post2Bic procedure
VarPost2Aic	Selected covariate (by group) given by Post2Bic procedure
VarPost2Gcv	Selected covariate (by group) given by Post2Bic procedure
group	Consecutive integers describing the grouping of the coefficients. Integers for penalized groups must run from 1 to the number of penalized groups
criterion	Factorial covariate indicating the model selection criterion used (BIC, AIC or GCV)

Author(s)

Vincent Thouvenot <vincentthouvenot4@gmail.com>

References

A.Antoniadis, Y.Goude, J-M. Poggi and V.Thouvenot. Selection de variables dans les mod<3><a8>les additifs avec des estimateurs en plusieurs <3><a9>tapes. Technical Report. <https://hal.archives-ouvertes.fr/hal-01116100>

V. Thouvenot, A. Pichavant, Y. Goude, A. Antoniadis, and J-M Poggi. Electricity forecasting using multi-step estimators of nonlinear additive models. Submitted, March 2015.

Annexe : Résultats complets des simulations présentées dans le chapitre 1

B.1 Tables récapitulant les performances en sélection de variables pour la simulation 1

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.18)	4(0)	4(0.29)	4(0.10)	4(0.25)	7 (1.26)	8(1.16)	4
MTZ	6(0)	6(0.18)	6(0)	6(0.29)	6(0.10)	6(0.25)	3 (1.26)	2(1.16)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.18)	0(0)	0(0.29)	0(0.10)	0(0.25)	3 (1.26)	4(1.16)	0

TABLE B.1 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 0.5 lorsque $t = 0$ (SNR=5.3)

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.36)	4(0)	4(0.20)	4(0.10)	4(0.12)	7 (1.23)	8 (1.12)	4
MTZ	6(0)	6(0.36)	6(0)	6(0.20)	6(0.10)	6(0.12)	3 (1.23)	2 (1.12)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.36)	0(0)	0(0.20)	0(0.10)	0(0.12)	3 (1.23)	4 (1.12)	0

TABLE B.2 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 1.5 lorsque $t = 0$ (SNR=3.1)

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.33)	4(0)	4(0.20)	4(0)	4(0)	7.5(1.28)	8 (1.22)	4
MTZ	6(0)	6(0.33)	6(0)	6(0.20)	6(0)	6(0)	2.5(1.28)	2 (1.22)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.33)	0(0)	0(0.20)	0(0)	0(0)	3.5(1.28)	4 (1.22)	0

TABLE B.3 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 3.5 lorsque $t = 0$ (SNR=2.0)

Lorsque les variables explicatives ne sont pas corrélées, les performances en termes de sélection de composantes des différentes méthodes issues du Group LASSO sont équivalentes, mis à part Post1Gcv, légèrement moins efficace.

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.35)	4(0)	4(0.24)	4(0.20)	4(0.52)	8(1.18)	8 (1.14)	4
MTZ	6(0)	6(0.35)	6(0)	6(0.24)	6(0.20)	6(0.52)	2(1.18)	2 (1.14)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.35)	0(0)	0(0.24)	0(0.20)	0(0.52)	4(1.18)	4 (1.14)	0

TABLE B.4 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 0.5 lorsque $t = 1$ (SNR=5.42)

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.38)	4(0)	4(0.20)	4(0.10)	4(0.22)	7(1.31)	8(1.24)	4
MTZ	6(0)	6(0.38)	6(0)	6(0.20)	6(0.10)	6(0.22)	3(1.31)	2(1.24)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.38)	0(0)	0(0.20)	0(0.10)	0(0.22)	3(1.31)	4(1.24)	0

TABLE B.5 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 1.5 lorsque $t = 1$ (SNR=3.1)

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.27)	4(0)	4(0.12)	4(0.18)	4(0.14)	7(1.14)	8(1.24)	4
MTZ	6(0)	6(0.27)	6(0)	6(0.12)	6(0.18)	6(0.14)	3(1.14)	2(1.24)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.27)	0(0)	0(0.12)	0(0.18)	0(0.14)	3(1.14)	4(1.24)	0

TABLE B.6 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 3.5 lorsque $t = 1$ (SNR=2.05)

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.25)	4(0)	4(0.14)	4(0.44)	4(0.39)	7(1.23)	8(1.23)	4
MTZ	6(0)	6(0.25)	6(0)	6(0.14)	6(0.44)	6(0.39)	3(1.23)	2(1.23)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.25)	0(0)	0(0.14)	0(0.44)	0(0.39)	3(1.23)	4(1.23)	0

TABLE B.7 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 0.5 lorsque $t = 2$ (SNR=5.38)

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	GAMSelect	GAMShrinkage	Idéal
MS	4(0)	4(0.25)	4(0)	4(0.14)	4(0.44)	4(0.39)	7(1.23)	8(1.23)	4
MTZ	6(0)	6(0.25)	6(0)	6(0.14)	6(0.44)	6(0.39)	3(1.23)	2(1.23)	6
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4(0)	4
MFP	0(0)	0(0.25)	0(0)	0(0.14)	0(0.44)	0(0.39)	3(1.23)	4(1.23)	0

TABLE B.8 – Capacité à sélectionner les composantes selon le critère utilisé avec une variance du bruit de 3.5 lorsque $t = 2$ (SNR=2.0)

B.2 Boîtes à moustaches des RMSE de la simulation 1

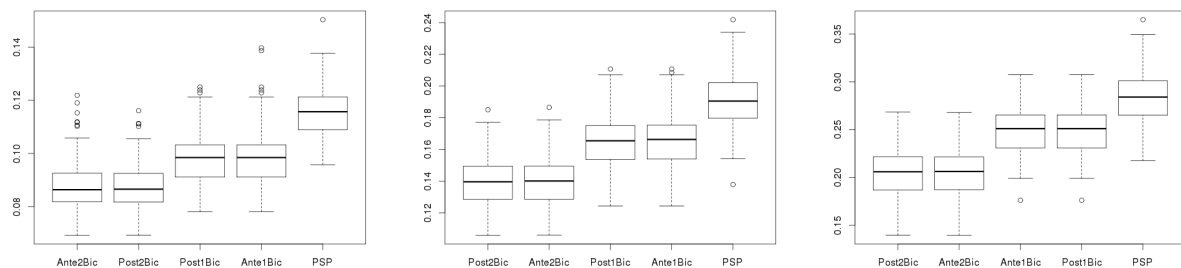


FIGURE B.1 – Boîtes à moustaches des RMSE lorsque la variance du bruit est respectivement 0.5, 1.5 et 3.5 et $t = 0$

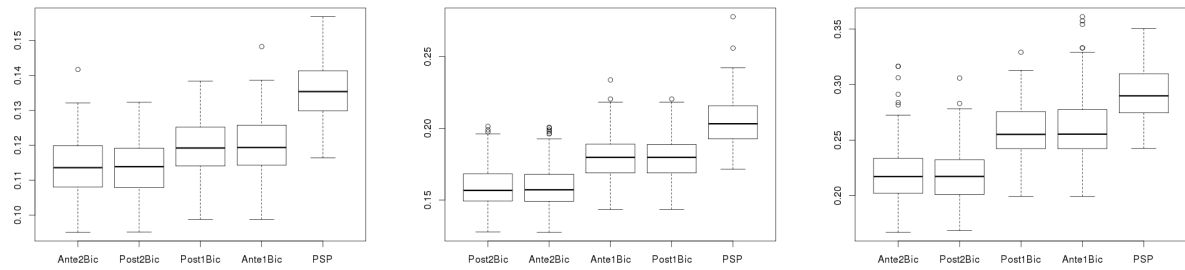


FIGURE B.2 – Boîtes à moustaches des RMSE lorsque la variance du bruit est respectivement 0.5, 1.5 et 3.5 et $t = 1$

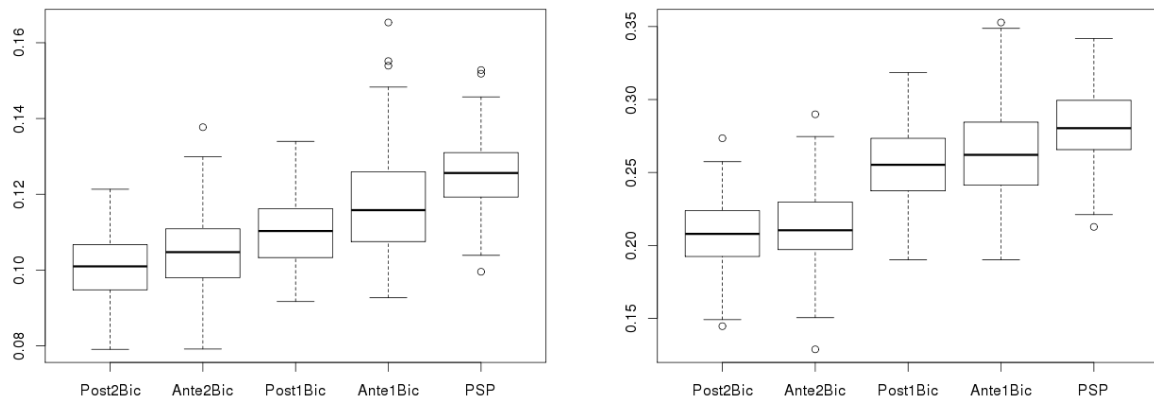


FIGURE B.3 – Boîtes à moustaches des RMSE lorsque la variance du bruit est respectivement 0.5 et 3.5 et $t = 2$

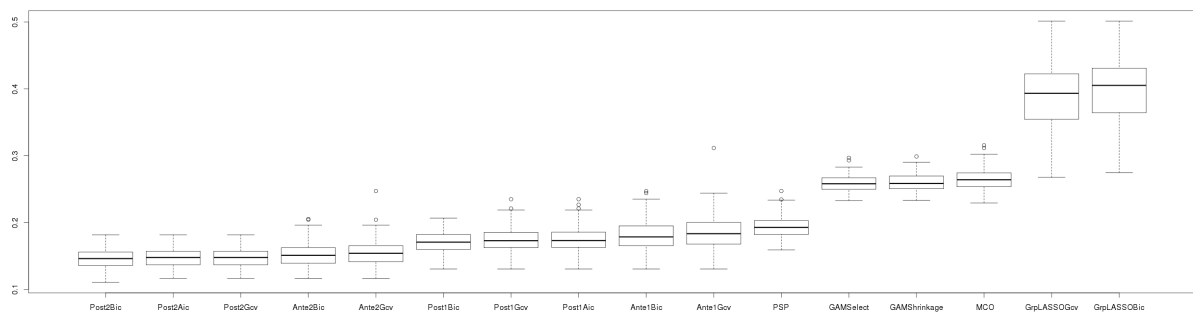


FIGURE B.4 – Boîtes à moustaches des RMSE lorsque la variance du bruit vaut 1.5 et $t = 2$ (SNR=3.1)

B.3 Phase de sélection de la pseudo-simulation

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	Idéal
MS	3(0)	3(0.24)	3(0)	3(0.17)	5(1.03)	4(0.68)	3
MTZ	11(0)	11(0.24)	11(0)	11(0.17)	9(1.03)	10(0.68)	11
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	3(0)	3(0)	3(0)	3(0)	3(0)	3(0)	3
MFP	0(0)	0(0.24)	0(0)	0(0.17)	2(1.03)	1(0.68)	0

TABLE B.9 – Capacité à sélectionner les composantes lorsque la variance du bruit est de 5.10^5 (SNR=4.40) pour la simulation type EDF 1

Critère	Post1Bic	Post1Gcv	Post2Bic	Post2Gcv	GrpLASSOBic	GrpLASSOGcv	Idéal
MS	3(0)	3(0.17)	3(0)	3(0.07)	5(1.03)	4(0.64)	3
MTZ	11(0)	11(0.17)	11(0)	11(0.07)	9(1.03)	10(0.64)	11
MFZ	0(0)	0(0)	0(0)	0(0)	0(0)	0(0)	0
MTP	3(0)	3(0)	3(0)	3(0)	3(0)	3(0)	3
MFP	0(0)	0(0.17)	0(0)	0(0.07)	2(1.03)	1(0.64)	0

TABLE B.10 – Capacité à sélectionner les composantes lorsque la variance du bruit est de 2.10^6 (SNR=2.20) pour la simulation type EDF 1



Annexe : Résultats complémentaires du chapitre 2

C.1 Lexique

Température nationale brute La température nationale est construite en utilisant un panier de 32 stations météorologiques. La France est divisée en plusieurs régions. Pour chaque région, des modèles de prévision de consommation sont estimés en testant des combinaisons différentes des températures des stations météorologiques de la zone considérée. Pour chaque région, la pondération minimisant un critère de performance (typiquement le RMSE) est conservée, puis ces poids ainsi obtenus sont pondérés par la consommation électrique de chaque région. Les poids actuellement utilisés ont été obtenus en utilisant les chroniques observées sur la période 1980-2010. La température peut être modélisée de manière physico-mathématique, en discrétisant des équations de la dynamique des fluides et en représentant des phénomènes physiques mis en jeu, ou statistiques, en ré-échantillonnant des observations et/ou simulant les principales caractéristiques statistiques des paramètres climatiques. En pratique, pour l'intégrer dans les modèles de prévision de consommation, elle est simulée de manière déterministe à un horizon journalier, probabiliste, en utilisant 51 scénarios, à un horizon hebdomadaire et bi-hebdomadaire et en utilisant 121 scénarios historiques à un horizon moyen terme.

Nébulosité La nébulosité, s'exprime en octa (de 0 pour un ciel entièrement découvert à 8 pour un ciel entièrement couvert), est observée par satellite et représente la couverture nuageuse. La nébulosité nationale est obtenue en pondérant les nébulosités de 32 stations météorologiques.

Vitesse du vent La vitesse du vent s'exprime en $m.s^{-1}$ et est calculée à 10 mètres du sol.

Projet LINKY Actuellement, les compteurs qui mesurent l'énergie consommée sont électromécaniques ou électroniques et nécessitent l'intervention de techniciens pour la mise en service, le relevé ou la modification de la puissance par exemple. Le compteur LINKY est un compteur qui peut recevoir et fournir des données automatiquement sans intervention humaine. Il est en interaction permanente avec le reste du réseau.

Données agrégées Données à un haut niveau d'agrégation comportant de nombreux clients avec une charge consommée élevée. Par exemple, la charge consommée sur le portefeuille EDF, en France ou en Europe.

Données désagrégées Données rassemblant un petit agrégat de consommateurs.

Données locales Données relativement désagrégées comportant plusieurs séries chronologiques qui ont un sens géographique.

Modèle long terme Modèle à un horizon long terme, c'est-à-dire permettant de prévoir plusieurs années de charges consommées à l'avance.

Modèle moyen terme Modèle à un horizon moyen terme, c'est-à-dire permettant de prévoir quelques mois ou quelques années de charges consommées.

Modèle court terme Modèle à un horizon court terme, c'est-à-dire permettant de prévoir la charge consommée dans une journée.

Correction court terme Correction d'un modèle moyen terme à un horizon de 24 heures grâce aux erreurs passées du modèle.

Modèle global Modèle estimé sans différenciation de l'heure de la journée.

Modèle par instant Un modèle différent est estimé par pas de temps journalier.

Effet climatisation Impact des températures chaudes sur la charge consommée.

Effet chauffage Impact des températures froides sur la charge consommée.

Modèle benchmark Modèle de référence.

C.2 Evaluation des performances en prévision des modèles

Un modèle (M1) a été estimé pour prévoir la charge consommée dans le futur. Supposons qu'il y ait n observations dans l'intervalle futur, que nous notons $\hat{Y}_t$ la prévision fournie par le modèle (M1) pour l'observation t et Y_t la charge consommée pour l'observation t . Nous nous intéressons à plusieurs critères de performance, qui sont donnés dans la Table C.1.

Critère	Formule
RMSE	$\sqrt{1/n \sum_{t=1}^n (\hat{Y}_t - Y_t)^2}$
MAPE	$1/n \sum_{t=1}^n \frac{ \hat{Y}_t - Y_t }{Y_t}$
MAE	$1/n \sum_{t=1}^n \hat{Y}_t - Y_t $

TABLE C.1 – Critère d'évaluation des performances

Le RMSE dépend de l'échelle de la variable à expliquer, il n'est donc utile que pour comparer des modèles d'une série unique. Il pénalise les fortes erreurs. Le MAE, qui dépend aussi de l'échelle de variable à expliquer, pénalise moins les fortes erreurs. Le MAPE est indépendant de l'échelle, mais il n'est pas symétrique. Si les valeurs de la variable à expliquer sont faibles ou si il y a des valeurs aberrantes, le MAPE peut être artificiellement fort.

C.3 Portefeuille EDF

C.3.1 Correction de la tendance sur le portefeuille EDF

La Figure C.1 présente la courbe de charge du portefeuille EDF sur la période étudiée non corrigée (gauche) et corrigée de la tendance (droite).

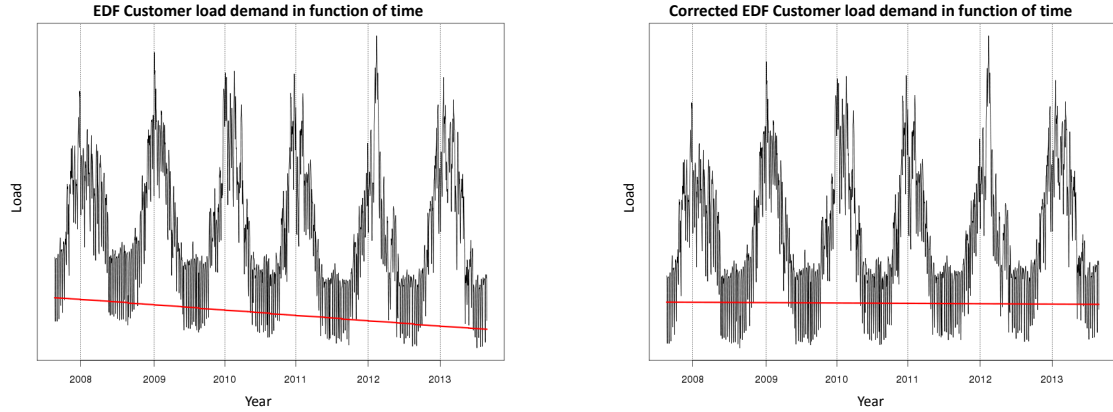


FIGURE C.1 – Charge consommée sur le portefeuille EDF non corrigée (gauche) et corrigée (droite) de la tendance

C.3.2 Post-traitement local

Nous homogénéisons les sous-ensembles de covariables en utilisant la dépendance entre les instants de la journée. Le post-traitement “local” consiste à soit ajouter une covariable non sélectionnée pour un instant donné lorsque celle-ci est présente pour les instants voisins, soit supprimer une variable sélectionnée lorsque celle-ci n’est pas sélectionnée pour les instants voisins. Pour ce type de traitement, nous avons choisi les deux règles suivantes :

- Pour une covariable non sélectionnée à un instant donné t , si elle est sélectionnée pour les 3 instants précédents ou suivants, ou pour les 2 instants encadrant l’instant t plus l’instant $t - 2$ ou $t + 2$, alors ajouter la covariable en question dans le sous-ensemble de variables associé à l’instant t .
- Pour une covariable sélectionnée à un instant donné t , si celle-ci ne l’est pas pour les deux instants précédents, encadrant ou suivant t , la supprimer du sous-ensemble de variables associé à l’instant t .

Les Figures C.2 et C.3 récapitulent les deux étapes du nettoyage.

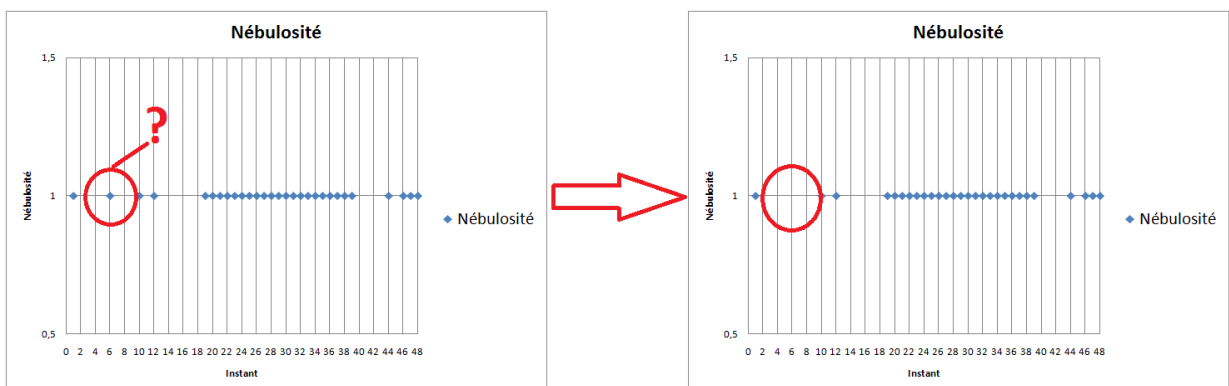


FIGURE C.2 – Post-traitement : suppression de covariables

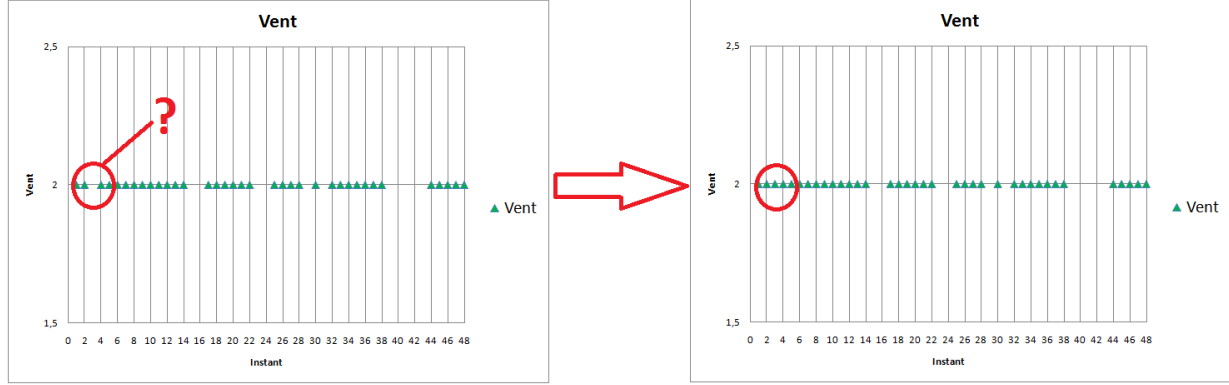


FIGURE C.3 – Post-traitement : ajout de covariables

C.3.3 Modèle Post2Bic>0.2

Le modèle Post2Bic>0.2 est donné dans l'équation (C.1).

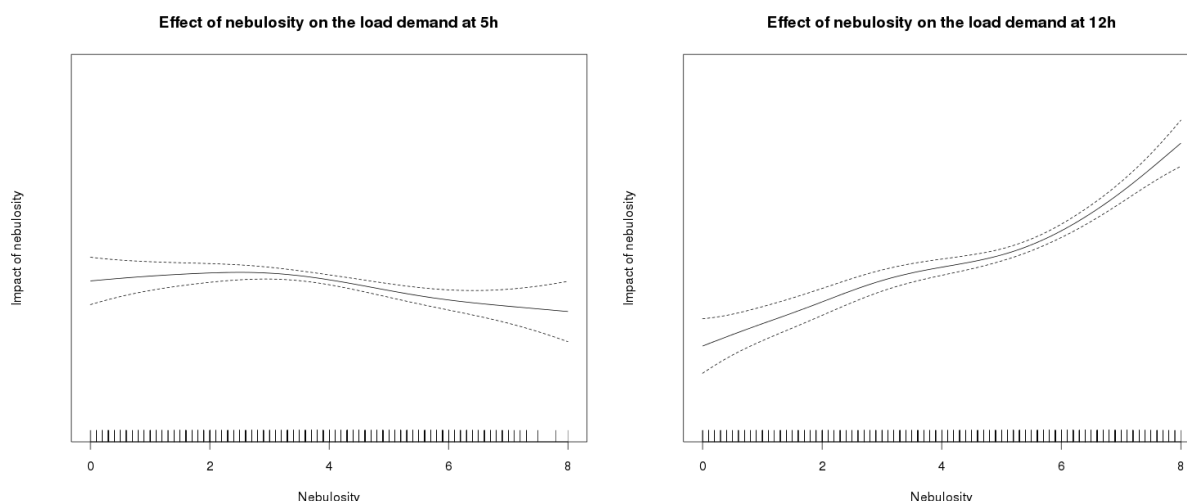
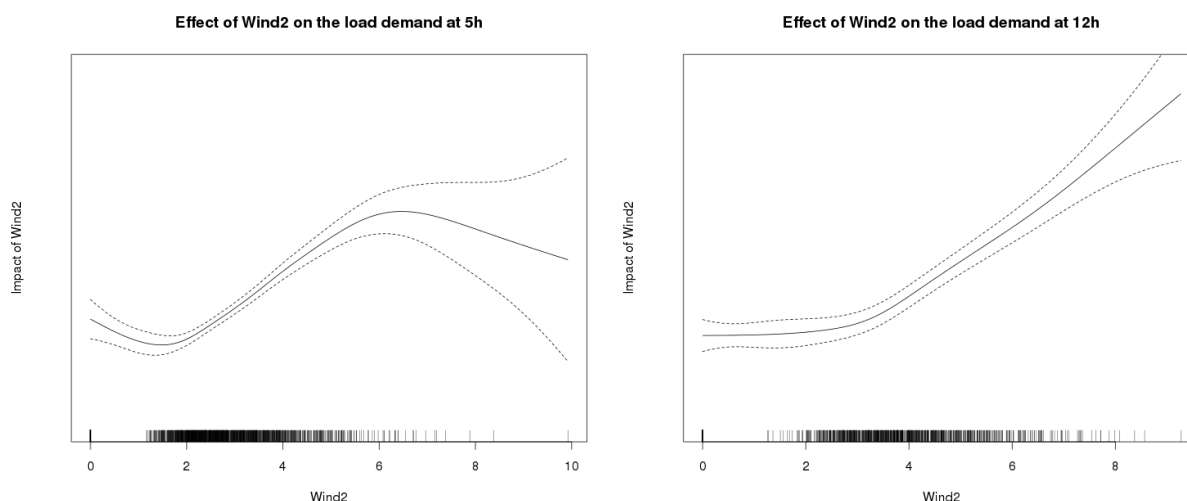
$$\begin{aligned}
 Y_t = & \beta_0 + \sum_{i=1}^7 \sum_{j=1}^3 \alpha_{ij} \mathbb{1}_{DayType_t=i} \mathbb{1}_{Offset_t=j} + \sum_{i=1}^{12} \beta_i T_t^{lag} \mathbb{1}_{Month_t=i} + s_1(Toy_t) \\
 & + g_1(CC_t) + g_2(T_t) + g_3(Wind2_t) + g_4(T_t^{lag,1}) \\
 & + k_1(\theta_t^{0.98}) + k_2(\theta_t^{0.982}) + k_3(\theta_t^{0.984}) + k_4(\theta_t^{0.990}) + k_5(\theta_t^{0.992}) + k_6(\theta_t^{0.994}) \\
 & + h_1(\theta_t^{Max,0.98}) + h_2(\theta_t^{Max,0.988}) \\
 & + l_1(\theta_t^{Min,0.98}) + l_2(\theta_t^{Min,0.984}) \\
 & + m_1(T_t^{Moy}) + m_2(T_t^{MoyJ1}) + m_3(T_t^{Max}) + m_4(T_t^{MaxJ1}) + \varepsilon_t,
 \end{aligned} \tag{C.1}$$

où

- $Wind2_t = W_t \mathbb{1}_{T_t \leq 15}$.
- $T_t^{lag,1}$ température retardée d'une heure
- θ_t^α température lissée, de coefficient α
- $\theta_t^{Max,\alpha}$ température maximale journalière lissée de coefficients α
- $\theta_t^{Min,\alpha}$ température minimale journalière lissée de coefficients α
- $T_t^{Moy}, T_t^{MoyJ1}, T_t^{Max}$ et T_t^{MaxJ1} respectivement la température moyenne journalière, la température moyenne de la veille, la température maximale journalière et la température maximale de la veille

La variable $Wind2$ a été intégrée dans le dictionnaire de covariables car le vent influence la température ressentie et influence donc le niveau de chauffage. Au dessus d'un certain seuil de température, le vent peut ne plus avoir d'effet sur la charge consommée. La nébulosité influence non seulement la température ressentie, mais aussi l'éclairage, ce qui explique que la nébulosité tronquée pour les températures élevées ne soit pas sélectionnée. La température retardée d'une heure est classiquement utilisée dans les modèles de prévision de consommation d'électricité et permet de tenir compte de l'inertie des bâtiments. Les températures moyennes du jour ou de la veille modélisent aussi cette inertie et lisse le signal de température. La dépendance de la charge consommée avec l'intensité des températures est modélisée par la présence des températures maximales du jour et de la veille.

Nous traçons les effets estimés à 5h et à midi de la nébulosité (Figure C.4) et du vent tronqué (Figure C.5). Sur ces figures, la même échelle a été conservée entre les graphiques de gauche et de droite. La charge consommée est plus faible à 5h que le midi, les effets sont donc lissés la nuit.

FIGURE C.4 – Effet estimé par $\text{Post2Bic} > 0.2$ de la nébulosité à 5h et à midiFIGURE C.5 – Effet estimé par $\text{Post2Bic} > 0.2$ du vent tronqué à 5h et à midi

Cependant, la forme de l'effet estimé de la nébulosité est différente pour ces deux instants. A midi, la demande d'électricité est croissante avec la nébulosité. Par contre, à 5h, l'effet de la nébulosité, qui est de plus faible amplitude, est décroissant, ce qui est cohérent car les nuits d'hiver avec une nébulosité faible ont tendance à être plus froides. L'effet estimé par $\text{Post2Bic} > 0.2$ du vent tronqué est relativement similaire à 5h et à 12h, excepté pour les valeurs extrêmes de vent. Lorsqu'il est faible, le vent tronqué a un effet relativement constant sur la charge consommée. À partir d'un seuil donné, son effet croît rapidement.

C.3.4 Performances et erreurs en prévision

La Figure C.6 présente le rapport entre les valeurs absolues ordonnées des erreurs de $\text{Post2Bic} > 0.2$ et Post2Gcv avec celles du modèle EDF. Seules les 70% plus élevées sont représentées par souci de lisibilité. Les erreurs fortes du modèle EDF ont tendance à être plus fortes que celles des deux autres modèles relativement aux erreurs faibles.

La Figure C.7 représente les erreurs du modèle Post2Gcv en estimation et en prévision. Les droites rouges représentent les droites constantes passant par 0. Les erreurs sont corrélées. La

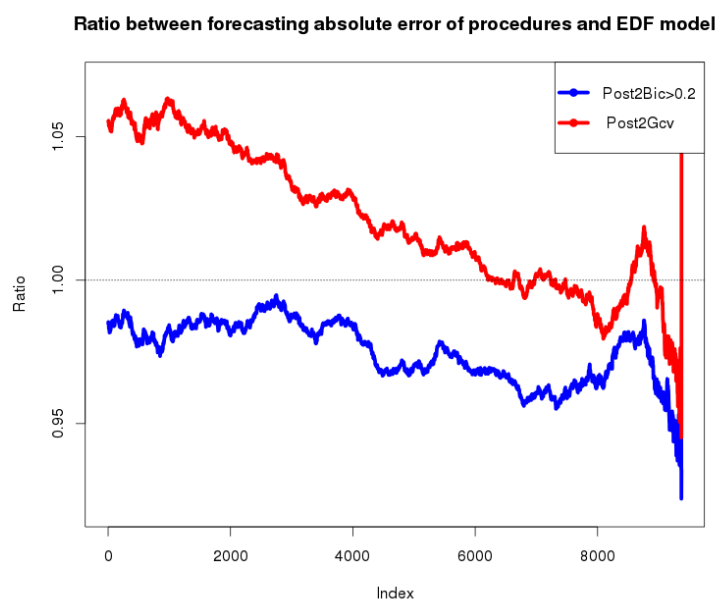


FIGURE C.6 – Ratio entre les valeurs absolues ordonnées des erreurs de $\text{Post2Bic} > 0.2$ et Post2Gcv avec celles du modèle EDF

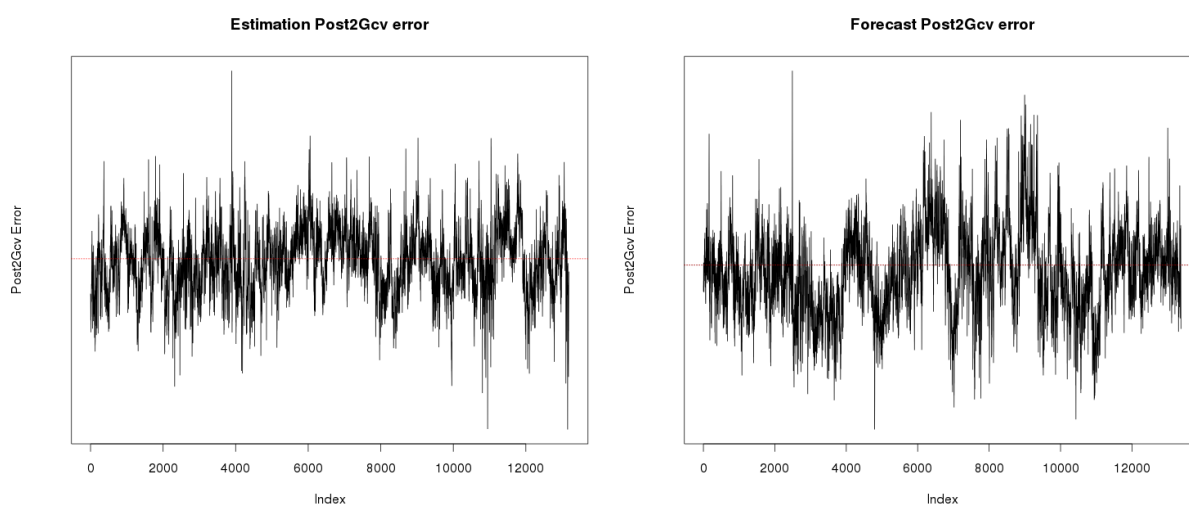


FIGURE C.7 – Erreurs des modèles en estimation et en prévision

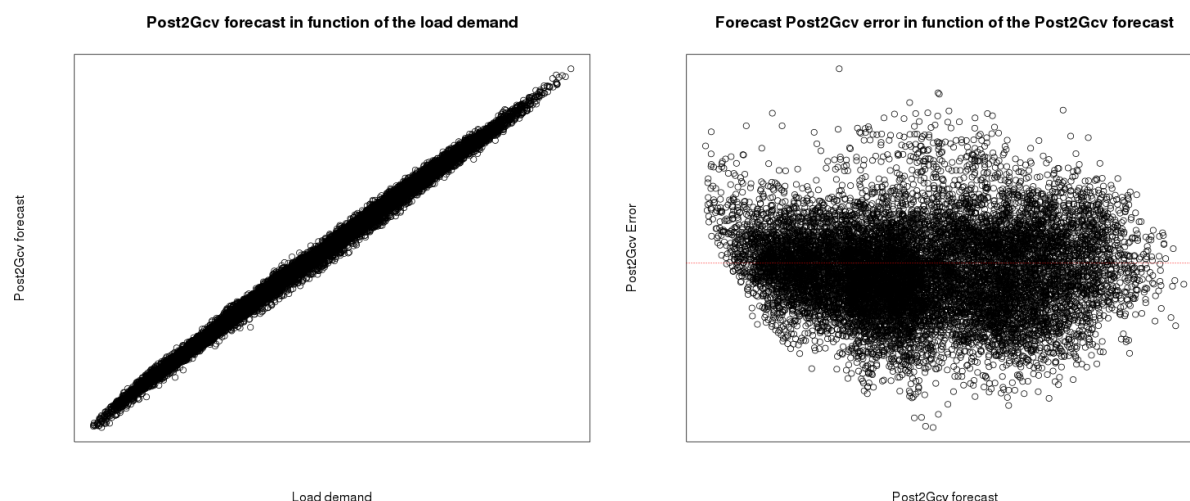


FIGURE C.8 – Etude des prévisions de Post2Gcv

Variable	R Carré	Corrélation
Température	0.002	-0.04
Température lissée expert	0.000	-0.01
Température moyenne	0.000	-0.02
Vent	0.016	0.12
Nébulosité	0.021	0.14
Toy	0.015	0.12

TABLE C.2 – Influence des covariables exogènes sur les erreurs commises en prévision par Post2Gcv

charge est sur-évaluée (erreurs négatives) en novembre (autour de la 3500ème observation) et sous-évaluée en avril (autour de 9000ème observation). Il y a eu un épisode très froid en avril 2013 comparé aux autres mois d'avril de l'échantillon. Près de 8 % des températures de ce mois sont inférieures à 5 degrés Celsius en 2013, alors qu'il y en a un peu moins de 3% pour les autres mois d'avril. Ces températures froides ont été nombreuses aux environs de 8-9h le matin, et influencent donc la charge consommée le matin. Il y a un pic d'erreurs classique autour cette heure en hiver. Le mois de novembre 2012 n'a pas eu d'épisodes extrêmes de températures, l'origine des erreurs plus élevées durant ce mois est donc différent.

La Figure C.8 donne les erreurs en prévision en fonction de la charge prédite et la charge prédite en fonction de la charge consommée. La charge prédite est bien corrélée avec la charge consommée et les valeurs faibles de la charge consommée sont sous-évaluées.

Nous nous intéressons sur la Figure C.9 aux erreurs de Post2Gcv en fonction de la température, de la température lissée de l'expert, de la température moyenne journalière, du vent, de la nébulosité et du moment dans l'année. Nous nous intéressons aussi aux modèles linéaires simples expliquant les erreurs de Post2Gcv par chaque covariable. Nous donnons dans la Table C.2 au R^2 de chaque modèle et la corrélation entre les erreurs et ces variables. Les variables exogènes ont peu d'impact sur les erreurs. La nébulosité les influence plus que la température (sous ses formes brute ou modifiées) ou le vent. Le vent comme la température modélisent principalement la thermosensibilité qui est modélisée par beaucoup de variables qui permettent plus de flexibilité, alors que l'effet de l'éclairage n'est principalement modélisé que par la nébulosité. Classiquement, l'éclairage est une cause importante de la consommation entre 19h et 22h en hiver.

Malgré cela, la température influence fortement les erreurs commises en avril 2013. Le modèle linéaire simple expliquant l'erreur en prévision commise au cours de ce mois par la température a

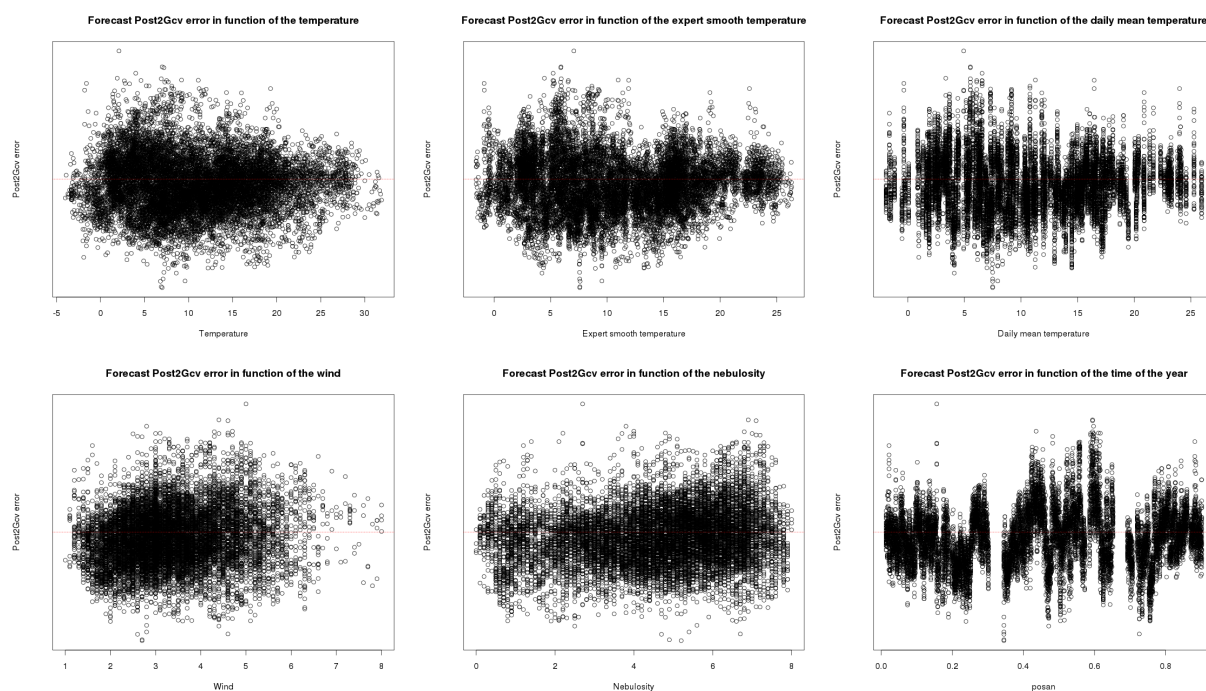


FIGURE C.9 – Erreurs de Post2Gcv en fonction de variables exogènes

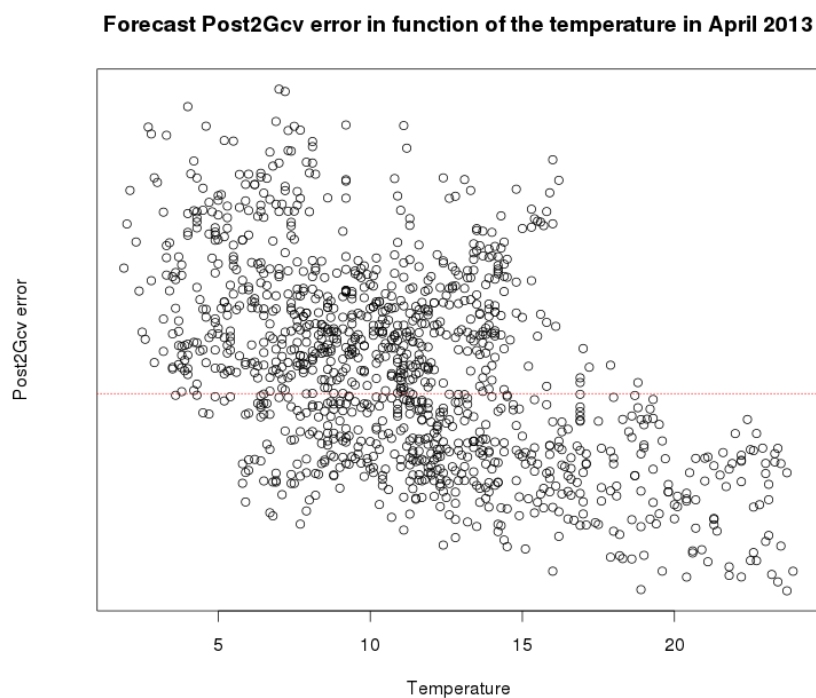


FIGURE C.10 – Erreurs commises par Post2Gcv en fonction de la température durant le mois d'avril 2013

un R^2 de 0.27. La corrélation entre la température et l'erreur est de -0.52. La Figure C.10 présente les erreurs commises par Post2Gcv en fonction de la température durant ce mois. Il y a une corrélation entre températures (notamment froides) et sous-évaluation de la charge consommée. Si nous raisonnons de même avec le mois de novembre, nous obtenons un R^2 de 0.008 et une corrélation de -0.09. Au cours de ce mois, les erreurs sont dues à une rupture de la consommation.

C.4 GEFCOM 2012

C.4.1 Performance moyen terme pour toutes les zones

Zone	Saturated	Post2Bic	Post2Aic	Post2Gcv	AggStation	OneStation
Zone 1	1753	1892	1748	1749	1896	1896
Zone 2	8854	9135	8959	9039	9288	9559
Zone 3	9554	9855	9666	9752	10023	10314
Zone 4	42	40	40	40	46	46
Zone 5	618	636	621	620	670	670
Zone 6	9014	9344	9142	9207	9529	9897
Zone 8	322	337	330	330	337	337
Zone 11	7913	8623	8376	8378	8773	8773
Zone 12	7958	8080	7858	7862	7960	8295
Zone 13	1815	1764	1745	1748	1764	1764
Zone 14	2273	2257	2207	2213	2246	2260
Zone 15	6142	5903	5873	5873	5941	5941
Zone 16	2667	2671	2570	2581	2669	2849
Zone 17	2172	2342	2173	2173	2175	2291
Zone 18	16079	16822	16270	16343	17191	18200
Zone 19	8050	7922	7958	7956	8117	8049
Zone 20	5449	5465	5437	5431	5418	5543

TABLE C.3 – RMSE pour les modèles MT

C.4.2 Court terme

La Figure C.11 présente les rapports entre les MAPE court et moyen termes pour toutes les zones de la compétition GEFCOM 2012. La Figure C.12, présente le rapport entre le MAPE de Post2Aic moyen terme et celui du modèle saturé moyen terme, et ce même rapport après correction court terme. La Figure C.13 présente le nombre de zones pour lesquelles un retard a été sélectionné. Nous n'avons représenté que les retards présents pour au moins quatre zones par souci de lisibilité. La Figure C.14 donne la boîte à moustaches de la dimension des modèles de correction court terme sélectionnés. Les modèles Post2Aic corrigés par (2.6) et par 2.7 sont appelés ST2 Post2Aic et ST Post2Aic respectivement. La Figure C.15 représente le rapport entre les MAPE en prévision de ST Post2Aic et ST2 Post2Aic.

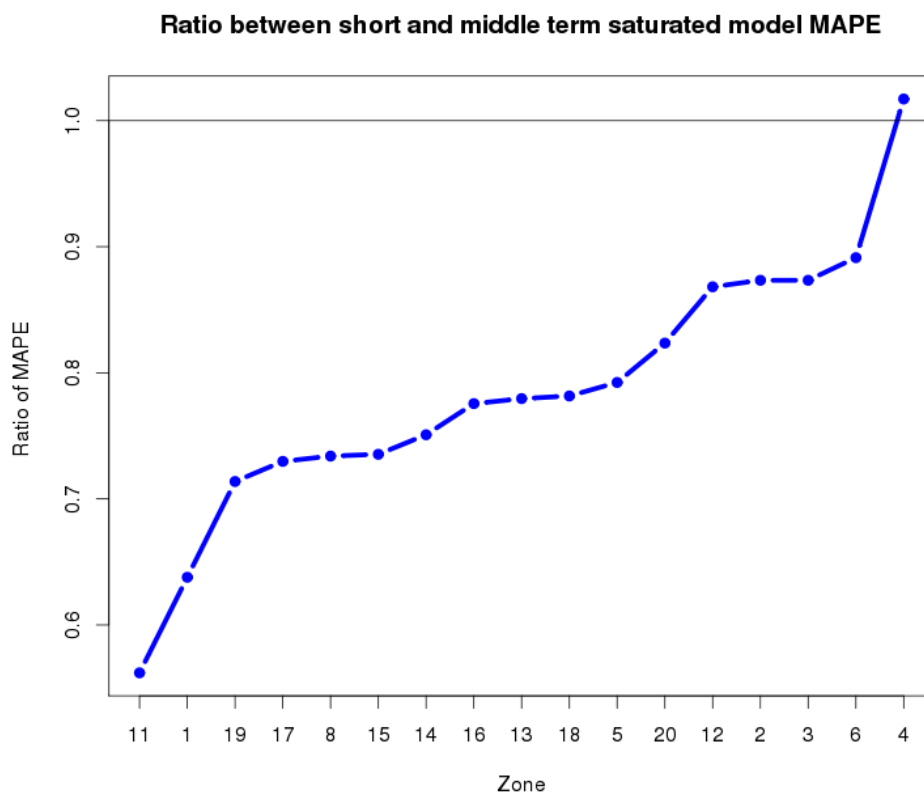


FIGURE C.11 – Rapport entre les MAPE court et moyen termes pour le modèle saturé pour GEFCom 2012

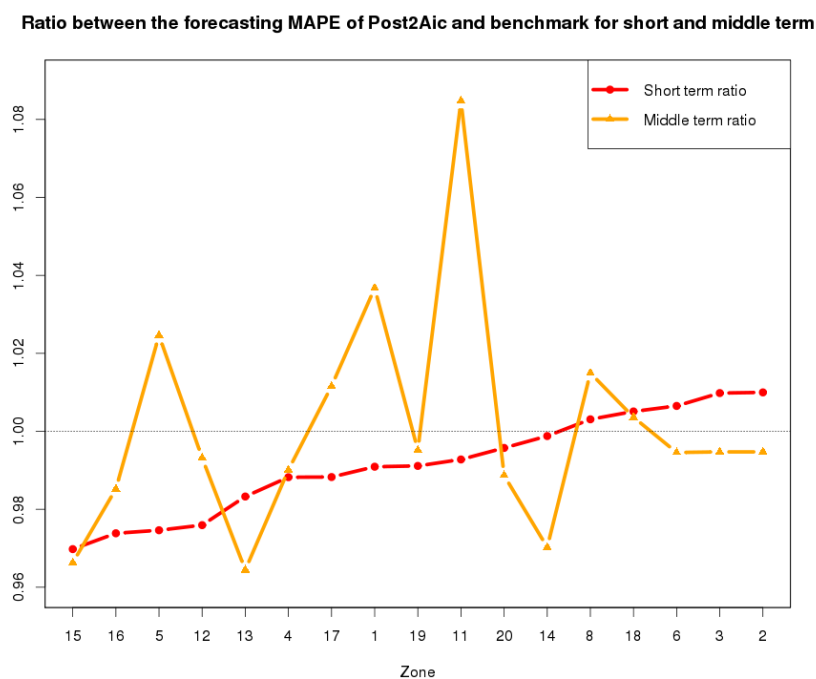


FIGURE C.12 – Comparaison des rapports de MAPE en prévision de Post2Aic et du modèle saturé à court et moyen termes

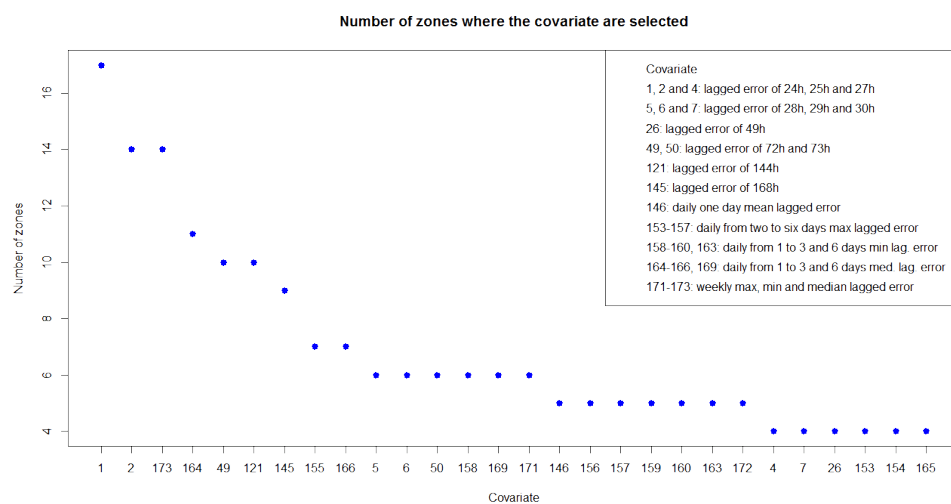


FIGURE C.13 – Nombre de zones où les retards sont sélectionnés

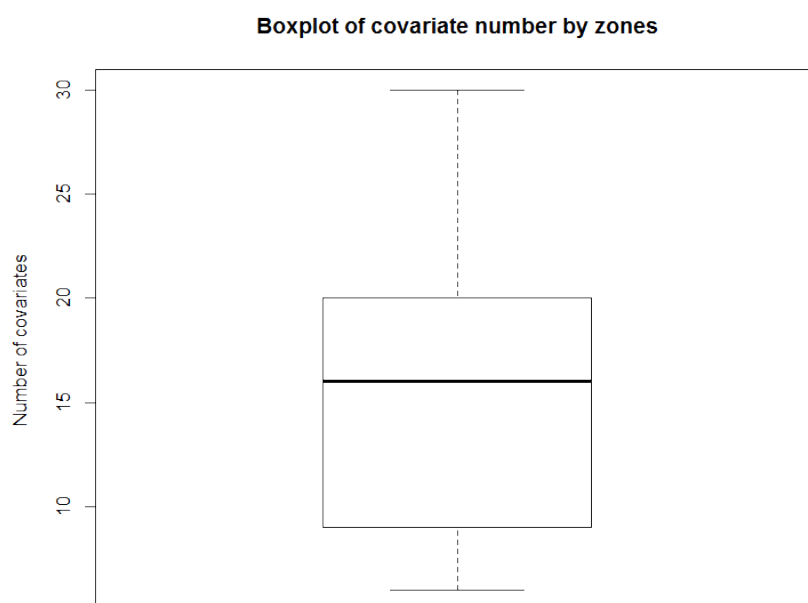


FIGURE C.14 – Nombre de retards sélectionnés par zone

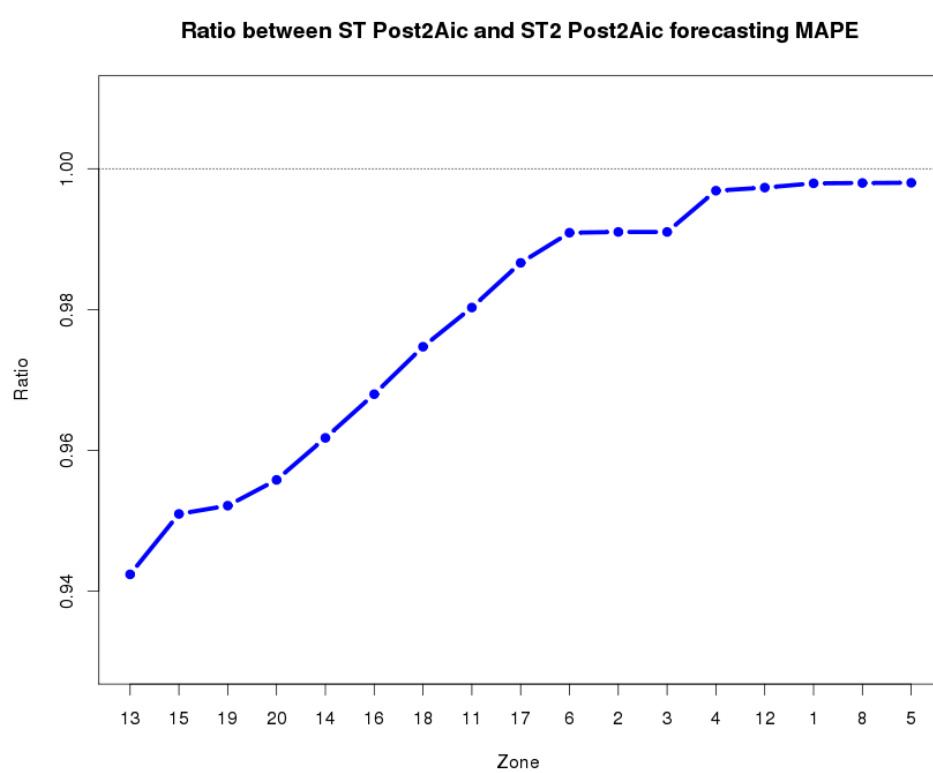


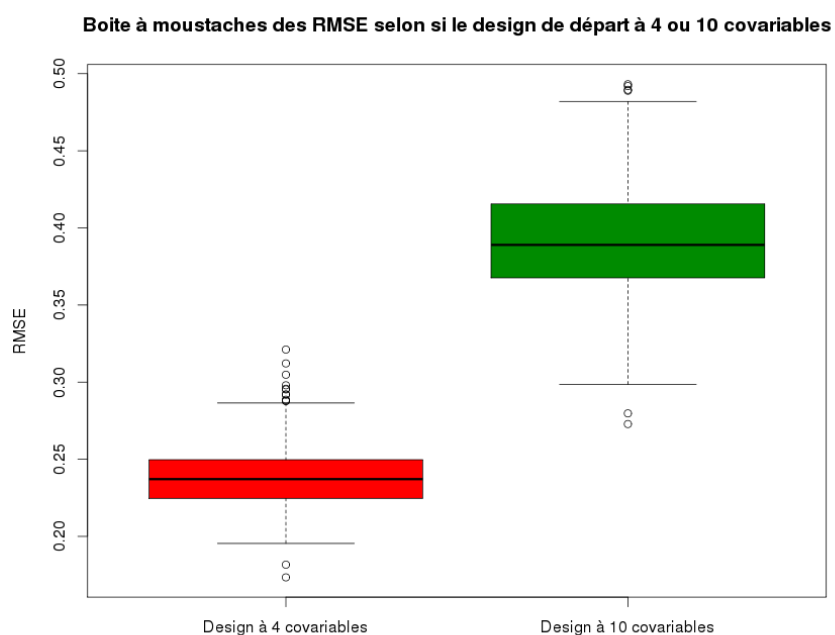
FIGURE C.15 – Ratio entre le MAPE de ST Post2Aic et ST2 Post2Aic

Annexe : Simulation associée au chapitre 3

Revenons à la simulation décrite dans la section 1.4 du chapitre 1 en considérant $t = 2$, un $SNR = 3.1$ et 3000 observations. Nous choisissons le modèle qui minimise le BIC. Il n'y a pas de faux négatif sur 500 répétitions de la simulation. Par contre, il y a des faux positifs (Table D.1). Dans un tiers des cas environ, le vrai design n'a pas été sélectionné. En plus de la possibilité d'introduire des faux positifs, en pratique, cette double pénalisation conduit à des résultats peu satisfaisants. L'estimateur résultant est moins performant que l'estimateur oracle construit lorsque S^* est supposé connu, et ce même si le vrai modèle est sélectionné. Pour l'illustrer, nous conservons uniquement les simulations où le vrai modèle a été sélectionné et estimons pour les 327 simulations concernées l'estimateur obtenu lorsque S^* est supposé connu. La Figure D.1 donne les boîtes à moustaches en prédiction lorsque le design de départ est S^* et lorsque le design de départ contient les dix covariables, sachant que le vrai modèle est sélectionné. La Figure justifie la séparation des phases de sélection et d'estimation. Une telle différence de performance s'explique par la valeur du paramètre de lissage de Group LASSO sélectionné : lorsqu'il y a dix covariables candidates, et donc des faux négatifs à supprimer, le BIC tend à sélectionner un paramètre de lissage plus fort pour exclure les covariables non influentes. Le biais introduit dans l'estimation des effets influents est alors plus fort, expliquant cette perte de qualité de prédiction. Le paramètre de pénalisation sélectionné est en moyenne de 0.03 avec un écart type de 0.003 dans le cas d'un design de départ de dix covariables et de 0.0001 lorsque S^* est utilisé, avec un écart type proche de 0.

Nombre de faux positifs	0	1	2	3
Nombre de simulation	327	126	42	5

TABLE D.1 – Nombre de faux positifs pour la simulation fondée sur [89]

FIGURE D.1 – RMSE en prédiction lorsque la pénalisation de Meier *et al.* [89] est utilisée directement sur le vrai design et sur le design à 10 covariables



Bibliographie

- [1] H. Akaike. A new look at the statistical model identification. *Automatic Control, IEEE Transactions on*, 19(6) :716–723, Dec 1974.
- [2] A.S. Al-Soliman and A.M. Al-Kandari. *Electric Load Modeling - Modeling and Model Construction*. Elsevier, Boston, 2010.
- [3] A. Antoniadis, X. Brossat, J. Cugliari, and J-M. Poggi. Prédiction d'un processus à valeurs fonctionnelles en présence de non stationnarités. Application à la consommation d'électricité. *Journal de la Société Française de Statistique*, 153(2) :52–78, 2012.
- [4] A. Antoniadis, X. Brossat, J. Cugliari, and J-M. Poggi. Clustering fonctionnel data using wavelets. *International Journal of Wavelets, Multiresolution and Information Processing*, 11(01) :1350003, 2013.
- [5] A. Antoniadis, X. Brossat, J. Cugliari, and J-M. Poggi. Une approche fonctionnelle pour la prédiction non-paramétrique de la consommation d'électricité. *Journal de la Société Française de Statistique*, 155(2) :202–219, 2014.
- [6] A. Antoniadis, X. Brossat, Y. Goude, J-M. Poggi, and V. Thouvenot. Automatic component selection in additive modeling of french national electricity load forecasting. *Submitted*, April 2015.
- [7] A. Antoniadis and J. Fan. Regularization of wavelet approximations. *Journal of the American Statistical Association*, 96 :939–967, 2001.
- [8] A. Antoniadis, I. Gijbels, and S. Lambert-Lacroix. Penalized estimation in additive varying coefficient models using grouped regularization. *Statistical Papers*, 55(3) :727–750, 2014.
- [9] A. Antoniadis, I. Gijbels, and A. Verhasselt. Variable Selection in Additive Models Using P-Splines. *Technometrics*, 54(4) :425–438, 2012.
- [10] A. Antoniadis, Y. Goude, J-M. Poggi, and V. Thouvenot. Sélection de variables dans les modèles additifs avec des estimateurs en plusieurs étapes. Technical report, 2015. <https://hal.archives-ouvertes.fr/hal-01116100>.
- [11] A. Antoniadis, E. Paparoditis, and T. Sapatinas. A functional wavelet–kernel approach for time series prediction. *Journal of the Royal Statistical Society : Series B*, 68(5) :837–857, 2006.
- [12] A. Ba, M. Sinn, Y. Goude, and P. Pompey. Adaptive learning of smoothing functions : application to electricity load forecasting. In *Advances in neural information processing systems*, pages 2510–2518, 2012.
- [13] F.R. Bach. Consistency of the group lasso and multiple kernel learning. *J. Mach. Learn. Res.*, 9 :1179–1225, June 2008.
- [14] I. Barrio, I. Arostegui, J. Quintana, and IRYSS-COPD Group. Use of generalised additive models to categorise continuous variables in clinical prediction. *BMC Medical Research Methodology*, 13(1) :83, 2013.
- [15] G.A. Begg and G. Marteinsdottir. Environmental and stock effects on spatial distribution and abundance of mature cod gadus morhua. *MARINE ECOLOGY PROGRESS SERIES*, 229 :245–262, 2002.

- [16] A. Belloni and V. Chernozhukov. Least squares after model selection in high-dimensional sparse models. *Bernoulli*, 19(2) :521–547, 05 2013.
- [17] S. Ben Taieb and R.J. Hyndman. A gradient boosting approach to the kaggle load forecasting competition. *International Journal of Forecasting*, 30(2) :382 – 394, 2014.
- [18] A. Berard and G. Hebrail. Searching time series with hadoop in an electric power company. In *Proceedings of the 2nd International Workshop on Big Data, Streams and Heterogeneous Source Mining : Algorithms, Systems, Programming Models and Applications*, pages 15–22. ACM, 2013.
- [19] P. Breheny and J. Huang. Group descent algorithms for nonconvex penalized linear and logistic regression models with grouped predictors. *Statistics and Computing*, pages 1–15, 2013.
- [20] L. Breiman. Better subset regression using the nonnegative garrote. *Technometrics*, 37(4) :373–384, November 1995.
- [21] E. Bruhns, G. Deurveilher, and J-S. Roy. A non-linear regression model for mid-term load forecasting and improvements in seasonality. *The 15th Power Systems Computation Conference, Liege, Belgium*, 2005.
- [22] D. W. Bunn and E. D. Farmer. *Comparative Models for Electrical Load Forecasting*. John Wiley, New York, 1985.
- [23] K.P. Burnham and D.R. Anderson. *Model selection and multimodel inference : a practical information-theoretic approach*. Springer, 2 edition, 2002.
- [24] R. Campo and P. Ruiz. Adaptive weather-sensitive short-term load forecasting. *IEEE Transactions on Power Systems*, 3 :592–600, 1987.
- [25] E. Cantoni, J. Mills Flemming, and E. Ronchetti. Variable selection in additive models by nonnegative garrote. *Statistical Modelling*, 11(3) :237–252, March 2006.
- [26] B. J. Chen, M. W. Chang, and C.-J. Lin. Load forecasting using support vector machines : a study on eunite competition 2001. *IEEE Transaction on Power Systems*, 19 :1821–1830, 2004.
- [27] H. Chen, T. Ngoc Cong, W. Yang, C. Tan, Y. Li, and Y. Ding. Progress in electrical energy storage system : A critical review. *Progress in Natural Science*, 19(3) :291 – 312, 2009.
- [28] H. Cho, Y. Goude, X. Brossat, and Q. Yao. Modelling and forecasting daily electricity load via curve linear regression. In A. Antoniadis, J-M. Poggi, and X. Brossat, editors, *Modeling and Stochastic Learning for Forecasting in High Dimensions*, volume 217 of *Lecture Notes in Statistics*, pages 35–54. Springer International Publishing, 2015.
- [29] H. Cho, Y. Goude, Q. Yao, and X. Brossat. Modelling and forecasting daily electricity load curves : a hybrid approach. *Journal of the American Statistical Association*, 108(501) :7–21, 2013.
- [30] J.H. Chow, F.F. Wu, and J.A. Momoh, editors. *Applied Mathematics for Restructured Electric Power Systems*. Power Electronics and Power Systems. Springer US, 2005.
- [31] G. Claeskens, T. Krivobokova, and J. D. Opsomer. Asymptotic properties of penalized spline estimators. *Biometrika*, 96(3) :529–544, 2009.
- [32] J. Cugliari. *Prévision non paramétrique de processus à valeurs fonctionnelles : application à la consommation d’électricité*. PhD thesis, 2011. Thèse de doctorat de Mathématiques dirigée par Antoniadis, A. et Poggi, J-M., Paris 11, <http://www.theses.fr/2011PA112234/document>.
- [33] C. De Boor. *A practical guide to splines*. Springer-Verlag New York, 1978.
- [34] A. De Moliner. Estimation de synchrones de consommations électriques avec calage dynamique sur des données de relevés d’index asynchrones. *Journal de la Société Française de Statistique*, 156(1) :1–10, 2015.

- [35] M. Devaine, P. Gaillard, Y. Goude, and G. Stoltz. Forecasting electricity consumption by aggregating specialized experts - a review of the sequential aggregation of specialized experts, with an application to slovakian and french country-wide one-day-ahead (half-) hourly predictions. *Machine Learning*, 90(2) :231–260, 2013.
 - [36] F. Dominici, A. McDermott, S. L. Zeger, and J.M. Samet. On the use of generalized additive models in time-series studies of air pollution and health. *American Journal of Epidemiology*, 156(3) :193–203, 2002.
 - [37] V. Dordonnat. *State-space Modelling for High Frequency Data : Three Applications to French National Electricity Load*. PhD thesis, 2009. Thèse de doctorat de Mathématiques dirigée par Koopman, S.J. et Ooms, M., Amsterdam, <http://chercheurs.edf.com/fichiers/fckeditor/Commun/Innovation/theses/TheseDordonnat.pdf>.
 - [38] V. Dordonnat, S. J. Koopman, M. Ooms, A. Dessertaine, and J. Collet. An hourly periodic state space model for modelling french national electricity load. *International Journal of Forecasting*, 24 :566–587, 2008.
 - [39] V. Dordonnat, A. Pichavant, and A. Pierrot. Gefcom2014 probabilistic electric load forecasting using time series and semi-parametric regression models. *Submitted in International Journal of Forecasting*, 2015.
 - [40] B. Efron, T. Hastie, I. Johnstone, and R. Tibshirani. Least angle regression. *The Annals of Statistics*, 32(2) :407–499, 04 2004.
 - [41] M.A. Efronson. *Multiple Regression Analysis*. Mathematical Methods for Digital Computers. Wiley, 1960.
 - [42] P.H.C. Eilers and B.D. Marx. Flexible smoothing with B-splines and penalties. *Statistical Science*, 11(2) :89–121, 05 1996.
 - [43] M. El Anbari and A. Mkhadri. Penalized regression combining the L1 norm and a correlation based penalty. Research Report RR-6746, 2008. <https://hal.inria.fr/inria-00343635/file/RR-6746.pdf>.
 - [44] J. Fan and J. Jiang. Generalized likelihood ratio tests for additive models. *Journal American Statistical Association*, 100 :890–907, 2005.
 - [45] S. Fan and L. Chen. Short-term load forecasting based on an adaptive hybrid method. *IEEE Transactions on Power Systems*, 21(1) :395–401, 2006.
 - [46] S. Fan and R.J. Hyndman. Short-term load forecasting based on a semi-parametric additive model. *Power Systems, IEEE Transactions on*, 27(1) :134–141, Feb 2012.
 - [47] I. E. Frank and J. H. Friedman. A statistical view of some chemometrics regression tools. *Technometrics*, 35 :109–148, 1993.
 - [48] G.M. Furnival and Robert W. Wilson, Jr. Regressions by leaps and bounds. *Technometrics*, 42(1) :69–79, February 2000.
 - [49] P. Gaillard. *Contributions à l’agrégation séquentielle robuste d’experts : travaux sur l’erreur d’approximation et la prévision en loi. Applications à la prévision pour les marchés de l’énergie*. PhD thesis, 2015. Thèse de doctorat de Mathématiques dirigée par Stoltz G., Paris 11, <http://pierre.gaillard.me/doc/theseGaillard.pdf>.
 - [50] P. Gaillard and Y. Goude. Forecasting electricity consumption by aggregating experts ; how to design a good set of experts. In A. Antoniadis, J-M Poggi, and X. Brossat, editors, *Modeling and Stochastic Learning for Forecasting in High Dimensions*, volume 217 of *Lecture Notes in Statistics*, pages 95–115. Springer International Publishing, 2015.
 - [51] P. Gaillard, Y. Goude, and R. Nedellec. Semi-parametric models for gefcom2014 probabilistic electric load and electricity price forecasting. *Submitted in International Journal of Forecasting*, 2015.
-

- [52] D. Gervini. Free-knot spline smoothing for functional data. *Journal of the Royal Statistical Society : Series B (Statistical Methodology)*, 68(4) :671–687, 2006.
- [53] Y. Goude. *Mixing Individual Predictors when Breaks Occur : A Risk Approach*. Université de Paris-Sud. Département de Mathématiques, 2006.
- [54] Y. Goude. *Mélange de prédicteurs : application à la prévision de consommation d’électricité*. PhD thesis, 2008. Thèse de doctorat de Mathématiques dirigée par Dacunha-Castelle, D., Paris 11, <http://www.theses.fr/2008PA112027>.
- [55] Y. Goude, R. Nedellec, and N. Kong. Local short and middle term electricity load forecasting with semi-parametric additive models. *Smart Grid, IEEE Transactions on*, 5(1) :440–446, Jan 2014.
- [56] H. Hahn, S. Meyer-Nieberg, and S. Pickl. Electric load forecasting methods : Tools for decision making. *European Journal of Operational Research*, 199(3) :902 – 907, 2009.
- [57] T. J. Hastie and R. J. Tibshirani. *Generalized additive models*. London : Chapman & Hall, 1990.
- [58] T. J. Hastie, R.J. Tibshirani, and J.H. Friedman. *The elements of statistical learning : data mining, inference, and prediction*. Springer series in statistics. Springer, New York, 2009. Autres impressions : 2011 (corr.), 2013 (7e corr.).
- [59] L. Hatton, P. Charpentier, and E. Matzner-Løber. Statistical estimation of the residential baseline. *To appear in IEEE Transactions on Power Systems*, 2015.
- [60] A. Hinojosa and V.H. Hoese. Short-term load forecasting using fuzzy inductive reasoning and evolutionary algorithms. *IEEE Transactions on Power Systems*, 25 :565–574, 2010.
- [61] R. R. Hocking. A biometrics invited paper. The analysis and selection of variables in linear regression. *Biometrics*, 32(1) :pp. 1–49, 1976.
- [62] A. E. Hoerl and R. W. Kennard. Ridge Regression : Biased Estimation for Nonorthogonal Problems. *Technometrics*, 12(1) :55–67, February 1970.
- [63] T. Hong, P. Pinson, and S. Fan. Global energy forecasting competition 2012. *International Journal of Forecasting*, 30(2) :357 – 363, 2014.
- [64] T. Hong, P. Pinson, S. Fan, H. Zareipour, A. Troccoli, and R.J. Hyndman. Probabilistic energy forecasting : State-of-the-art 2015. *Submitted in International Journal of Forecasting*, 2015.
- [65] T. Hong, P. Wang, and L. White. Weather station selection for electric load forecasting. *International Journal of Forecasting*, 31(2) :286 – 295, 2015.
- [66] J. Huang, P. Breheny, and S. Ma. A selective review of group selection in high-dimensional models. *Statistical Science*, 27(4) :481–499, 11 2012.
- [67] J. Huang, J. L. Horowitz, and F. Wei. Variable selection in nonparametric additive models. *The Annals of Statistics*, 38(4) :2282–2313, 08 2010.
- [68] S. J. Huang and K. R. Shih. Short-term load forecasting via ARMA model identification including nongaussian process considerations. *IEEE Transactions on Power Systems*, 18(2) :673–679, 2003.
- [69] R.J. Hyndman and S. Fan. Density forecasting for long-term peak electricity demand. *Power Systems, IEEE Transactions on*, 25(2) :1142–1153, May 2010.
- [70] F.X. Jollois, J-M. Poggi, and B. Portier. Three non-linear statistical methods for analyzing PM10 pollution in Rouen area. *CSBIGS*, 3 :1–17, 2009.
- [71] K. Kato. Group Lasso for high dimensional sparse quantile regression models. *ArXiv e-prints*, March 2011. <http://adsabs.harvard.edu/abs/2011arXiv1103.1458K>.
- [72] K. Kato. Two-step estimation of high dimensional additive models. Technical report, jul 2012. <http://adsabs.harvard.edu/abs/2012arXiv1207.5313K>.

- [73] A. Khotanzad, R. Afkhami-Rohani, and D. Maratukulam. ANNSTLF-Artificial neural network short-term load forecaster generation three. *Power Systems, IEEE Transactions on*, 13(4) :1413–1422, Nov 1998.
- [74] R. Koenker and G. Bassett Jr. Regression quantiles. *Econometrica : Journal of the Econometric Society*, pages 33–50, 1978.
- [75] A.N. Kolmogorov. On tables of random numbers. *Theoretical Computer Science*, 207(2) :387 – 395, 1998.
- [76] S. Kullback and R. A. Leibler. On information and sufficiency. *Ann. Math. Statist.*, 22(1) :79–86, 03 1951.
- [77] T. Launay. *Bayesian methods for electricity load forecasting*. PhD thesis, 2012. Thèse de doctorat de Mathématiques dirigée par Philippe, A., Nantes, <https://tel.archives-ouvertes.fr/tel-00766237>.
- [78] T. Launay, A. Philippe, and S. Lamarche. Construction of an informative hierarchical prior distribution. Application to electricity load forecasting. <https://hal.archives-ouvertes.fr/hal-00625117>, September 2011.
- [79] T. Launay, A. Philippe, and S. Lamarche. On particle filters applied to electricity load forecasting. *arXiv :1210.0770*, 2012.
- [80] D.D. Lee and H.S. Seung. Algorithms for non-negative matrix factorization. In T.K. Leen, T.G. Dietterich, and V. Tresp, editors, *Advances in Neural Information Processing Systems 13*, pages 556–562. MIT Press, 2001.
- [81] H. Leeb and B. M. Potscher. Can one estimate the conditional distribution of post-model-selection estimators? *The Annals of Statistics*, 34(5) :2554–2591, 10 2006.
- [82] V. Lefieux. *Modèles semi-paramétriques appliquées à la prévision des séries temporelles. Cas de la consommation d’électricité*. PhD thesis, 2007. Thèse de doctorat de Mathématiques dirigée par Carbon, M. et Delecroix, M., Rennes 2, <https://tel.archives-ouvertes.fr/tel-00179866/document>.
- [83] Y. Lin and H.H. Zhang. Component selection and smoothing in multivariate nonparametric regression. *The Annals of Statistics*, 34(5) :2272–2297, Oct 2006.
- [84] J.R. Lloyd. Gefcom2012 hierarchical load forecasting : Gradient boosting machines and gaussian processes. *International Journal of Forecasting*, 30(2) :369–374, 2014.
- [85] C. L. Mallows. Some comments on Cp. *Technometrics*, 15 :661–675, 1973.
- [86] G. Marra and S. Wood. Practical variable selection for generalized additive models. *Computational Statistics & Data Analysis*, 55(7) :2372 – 2387, 2011.
- [87] P. Mathis. *Les énergies : comprendre les enjeux*. Quae, 2011.
- [88] L. Meier, S. Van De Geer, and P. Bühlmann. The group lasso for logistic regression. *Journal of the Royal Statistical Society, Series B*, 70 :53–71, 2008.
- [89] L. Meier, S. van de Geer, and P. Bühlmann. High-dimensional additive modeling. *The Annals of Statistics*, 37(6B) :3779–3821, Dec 2009.
- [90] N. Molinari, JF. Durand, and R. Sabatier. Bounded optimal knots for regression splines. *Computational Statistics & Data Analysis*, 45(2) :159 – 178, 2004.
- [91] R. Nedellec, J. Cugliari, and Y. Goude. Gefcom2012 : Electric load forecasting and backcasting with semi-parametric models. *International Journal of Forecasting*, 30(2) :375 – 381, 2014.
- [92] J. Nowicka-Zagrajeka and R. Weron. Modeling electricity loads in california : ARMA models with hyperbolic noise. *Signal Processing*, 82 :1903 –1915, 2002.
- [93] L. O’Brien and P. Rago. An application of the generalized additive model to groundfish survey data with atlantic cod off the northeast coast of the united states as an example. *NAFO Sci. Coun. Studies*, 28 :79–95, 1996.

- [94] A. Pierrot and Y. Goude. Short-term electricity load forecasting with generalized additive models. *Proceedings of ISAP power*, pages 593–600, 2011.
- [95] J. M. Poggi. Pr vision non param trique de la consommation  lectrique. *Revue de Statistique Appliqu e*, 42(4) :83–98, 1994.
- [96] P. Pompey, A. Bondu, Y. Goude, and M. Sinn. Massive-scale simulation of electrical load in smart grids using generalized additive models. In A. Antoniadis, J-M. Poggi, and X. Brossat, editors, *Modeling and Stochastic Learning for Forecasting in High Dimensions*, volume 217 of *Lecture Notes in Statistics*, pages 193–212. Springer International Publishing, 2015.
- [97] Z. Qin, K. Scheinberg, and D. Goldfarb. Efficient block-coordinate descent algorithms for the group lasso. *Mathematical Programming Computation*, 5(2) :143–169, 2013.
- [98] R. Ramanathan, R. Engle, C.W.J. Granger, F. Vahid-Araghi, and C. Brace. Short-run forecasts of electricity loads and peaks. *International Journal of Forecasting*, 13 :161–174, 1997.
- [99] J. Rissanen. Modeling by shortest data description. *Automatica*, 14(5) :465 – 471, 1978.
- [100] D. Ruppert. Selecting the number of knots for penalized splines. *Journal of Computational and Graphical Statistics*, 11(4) :735–757, December 2002.
- [101] C. Scherrer, T. Ambuj, M. Halappanavar, and D. Haglin. Feature clustering for accelerating parallel coordinate descent. In F. Pereira, C.J.C. Burges, L. Bottou, and K.Q. Weinberger, editors, *Advances in Neural Information Processing Systems 25*, pages 28–36. Curran Associates, Inc., 2012.
- [102] G. Schwarz. Estimating the dimension of a model. *The Annals of Statistics*, 6(2) :461–464, 1978.
- [103] C. E. Shannon. A mathematical theory of communication. *Bell System Technical Journal*, 27(3) :379–423, 1948.
- [104] R. Shibata. *Statistical Aspects of Model Selection*. Springer Berlin Heidelberg, 1989.
- [105] C. Singleton and N. Charlton. A refined parametric model for short term load forecasting. *International Journal of Forecasting*, 30(2) :364 – 368, April 2014.
- [106] C.J. Stone. Additive regression and other nonparametric models. *The Annals of Statistics*, 13(2) :689–705, 06 1985.
- [107] G. Swartzman, C. Huang, and S. Kaluzny. Spatial analysis of bering sea groundfish survey data using generalized additive models. *Canadian Journal of Fisheries and Aquatic Sciences*, 49(7) :1366–1378, 1992.
- [108] J. W. Taylor. Triple seasonal methods for short-term load forecasting. *European Journal of Operational Research*, 204 :139–152, 2010.
- [109] J.W. Taylor. Short-term load forecasting with exponentially weighted methods. *IEEE Transactions on Power Systems*, 27 :458–464, 2012.
- [110] RL. Thorndike. Who belongs in the family ? *Psychometrika*, 18(4) :267–276, 1953.
- [111] V. Thouvenot, A. Pichavant, Y. Goude, A. Antoniadis, and J-M Poggi. Electricity forecasting using multi-stage estimators of nonlinear additive models. *To appear in IEEE Transactions on Power Systems*, 2015.
- [112] R. Tibshirani. Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society, Series B*, 58 :267–288, 1994.
- [113] P. Tseng. Convergence of a block coordinate descent method for nondifferentiable minimization. *Journal of Optimization Theory and Applications*, 109(3) :475–494, 2001.
- [114] J. Ulbricht. *Variable Selection in Generalized Linear Models*. Verlag Dr. Hut, 2010.

- [115] G. Wahba and P. Craven. Smoothing noisy data with spline functions. Estimating the correct degree of smoothing by the method of generalized cross-validation. *Numerische Mathematik*, 31 :377–404, 1978.
 - [116] H. Wang, B. Li, and C. Leng. Shrinkage tuning parameter selection with a diverging number of parameters. *Journal of the Royal Statistical Society : Series B (Statistical Methodology)*, 71(3) :671–683, 2009.
 - [117] Wang, H. and Banerjee, A. . Randomized Block Coordinate for Online and Stochastic Optimization . *ArXiv e-prints*, 2014. <http://arxiv.org/pdf/1407.0107v3.pdf>.
 - [118] L. Wasserman. *All of Nonparametric Statistics (Springer Texts in Statistics)*. Springer-Verlag New York, Inc., Secaucus, NJ, USA, 2006.
 - [119] S. Wood. *Generalized Additive Models : An Introduction with R*. Chapman and Hall/CRC, 2006.
 - [120] S. Wood, Y. Goude, and S. Shaw. Generalized additive models for large data sets. *Journal of the Royal Statistical Society : Series C (Applied Statistics)*, 64(1) :139–155, 2015.
 - [121] C. Xia, J. Wang, and K. McMenemy. Short, medium and long term load forecasting model and virtual load forecaster based on radial basis function neural networks. *International Journal of Electrical Power and Energy Systems*, 32(7) :743 – 750, 2010.
 - [122] M. Yuan and Y. Lin. Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society, Series B*, 68 :49–67, 2006.
 - [123] S. Zhou and X. Shen. Spatially adaptive regression splines and accurate knot selection schemes. *Journal of the American Statistical Association*, 96(453) :247–259, 2001.
 - [124] H. Zou and T. Hastie. Regularization and variable selection via the elastic net. *Journal of the Royal Statistical Society, Series B*, 67 :301–320, 2005.
-

