

Sommaire

Introduction

Prévision de consommation électrique

Protection de données personnelles

Un monde de données

Conclusion

- ▶ **CPGE - Spécialité Maths Physique**
- ▶ ENSAI (Ecole Nationale de Statistique et de l'Analyse de l'Information) - Spécialité Génie Statistique - Rennes
- ▶ Master 2 de Statistique et d'Econométrie - Université de Rennes 1
- ▶ Thèse de Mathématiques Appliquées sous la direction d'Antoniadis, de Poggi et de Goude (tuteur industriel) - Université d'Orsay et EDF R&D
- ▶ Data Scientist - Thales Communication & Security

- ▶ CPGE - Spécialité Maths Physique
- ▶ ENSAI (Ecole Nationale de Statistique et de l'Analyse de l'Information) - Spécialité Génie Statistique - Rennes
- ▶ Master 2 de Statistique et d'Econométrie - Université de Rennes 1
- ▶ Thèse de Mathématiques Appliquées sous la direction d'Antoniadis, de Poggi et de Goude (tuteur industriel) - Université d'Orsay et EDF R&D
- ▶ Data Scientist - Thales Communication & Security

- ▶ CPGE - Spécialité Maths Physique
- ▶ ENSAI (Ecole Nationale de Statistique et de l'Analyse de l'Information) - Spécialité Génie Statistique - Rennes
- ▶ Master 2 de Statistique et d'Econométrie - Université de Rennes 1
- ▶ Thèse de Mathématiques Appliquées sous la direction d'Antoniadis, de Poggi et de Goude (tuteur industriel) - Université d'Orsay et EDF R&D
- ▶ Data Scientist - Thales Communication & Security

- ▶ CPGE - Spécialité Maths Physique
- ▶ ENSAI (Ecole Nationale de Statistique et de l'Analyse de l'Information) - Spécialité Génie Statistique - Rennes
- ▶ Master 2 de Statistique et d'Econométrie - Université de Rennes 1
- ▶ Thèse de Mathématiques Appliquées sous la direction d'Antoniadis, de Poggi et de Goude (tuteur industriel) - Université d'Orsay et EDF R&D
- ▶ Data Scientist - Thales Communication & Security

- ▶ CPGE - Spécialité Maths Physique
- ▶ ENSAI (Ecole Nationale de Statistique et de l'Analyse de l'Information) - Spécialité Génie Statistique - Rennes
- ▶ Master 2 de Statistique et d'Econométrie - Université de Rennes 1
- ▶ Thèse de Mathématiques Appliquées sous la direction d'Antoniadis, de Poggi et de Goude (tuteur industriel) - Université d'Orsay et EDF R&D
- ▶ Data Scientist - Thales Communication & Security

oooooooooooooooooooo

oooooooooooooooooooo

oooooooooooo

ENSAI

- ▶ Campus de Ker Lann proche de Rennes
- ▶ Habilitée par la CTI à délivrer le diplôme d'ingénieur conférant le grade de master
- ▶ Triple compétence statistique-économétrie-informatique
- ▶ Deux cycles
 - ▶ Fonctionnaire (INSEE et Ministères)
 - ▶ Ingénieur
- ▶ 6 filières de spécialisation pour la 3ème année des ingénieurs
 - ▶ Marketing Quantitatif et Revenu Management
 - ▶ Statistique et Ingénierie des données
 - ▶ Génie Statistique
 - ▶ Statistique pour les sciences de la vie
 - ▶ Gestion des risques et ingénierie financière
 - ▶ Ingénierie statistique des territoires et de la santé
- ▶ Recrutement à partir de bac + 2 : 50 % de MP, 15 % de B/L (économie et Sciences sociales) et ENS Cachan D2 (économie et Gestion), 35 % sur titres (DUT Stid, Licence de mathématique, étudiants internationaux, etc.)

Contribution des Statistiques en entreprises : illustration avec les filières de l'ENSAI

- ▶ Marketing Quantitatif et Revenu Management
 - ▶ Etude comportement de clients d'un grand fournisseur de jeu vidéo
 - ▶ Ciblage de clients dans le luxe
- ▶ Statistique et Ingénierie des données
 - ▶ Implémentation d'algorithmes passant à l'échelle pour de la cybersécurité
- ▶ Statistique pour les sciences de la vie
 - ▶ Etude de validation de médicaments
- ▶ Gestion des risques et ingénierie financière
 - ▶ Modélisation des défauts de paiement pour une banque
 - ▶ Scoring de clients pour une banque
- ▶ Ingénierie statistique des territoires et de la santé
 - ▶ Evaluation économique des produits de sante
- ▶ Génie Statistique
 - ▶ Maintenance prédictive pour l'industrie
 - ▶ Modélisation de la consommation électrique

Sommaire

Introduction

Prévision de consommation électrique

Contexte industriel

Modélisation statistique

Justification métier de la sélection de variables

Régression pénalisée

Portefeuille global d'un fournisseur

Application locale

Apport des Mathématiques

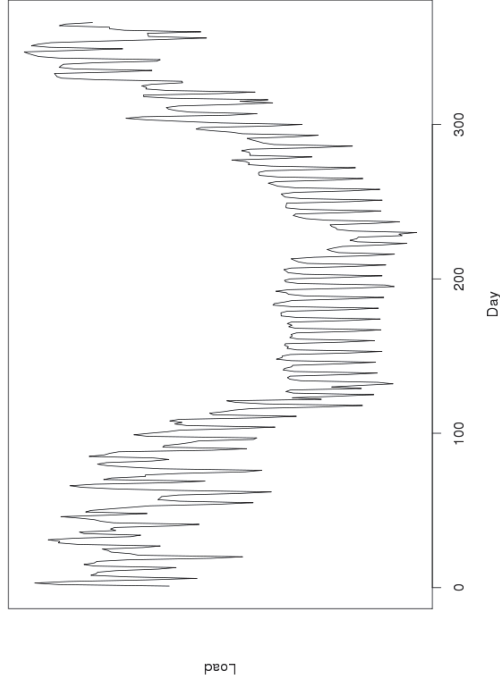
Protection de données personnelles

Un monde de données

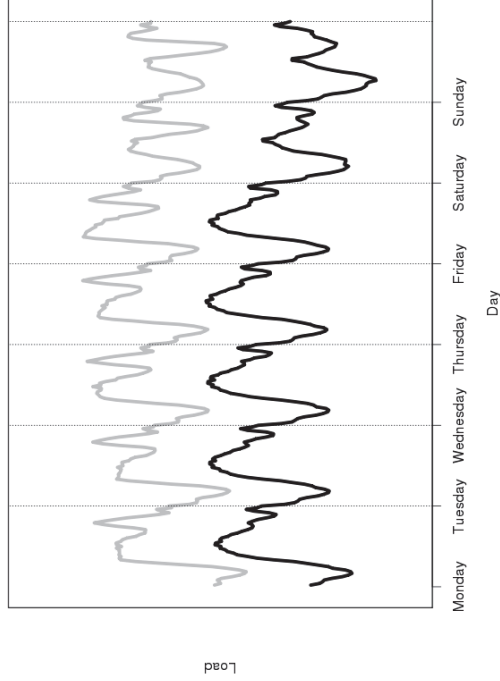
Conclusion

Courbe de charge française

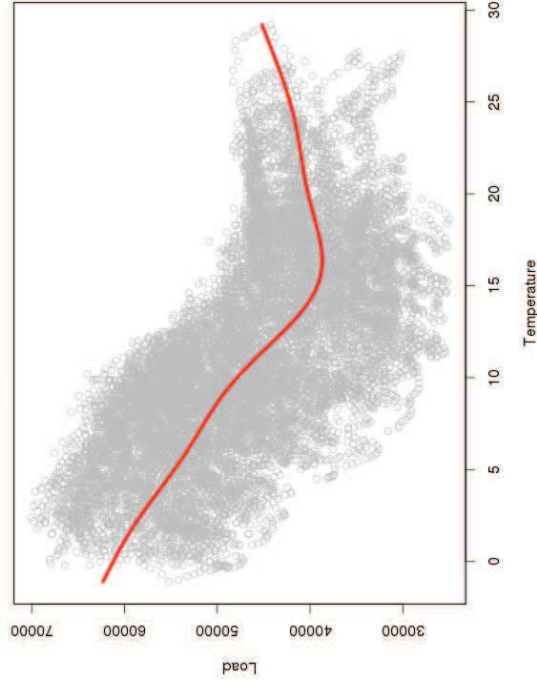
Daily load demand for 2008



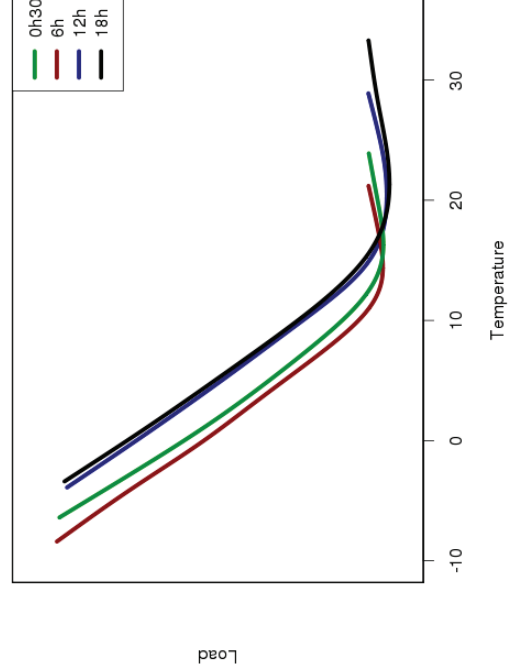
Load demand for one summer week (black) and one winter week (grey)



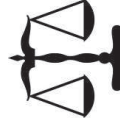
Effect of the temperature on the load demand in 2008



Effect of temperature along the hour



L'électricité ne se stocke pas



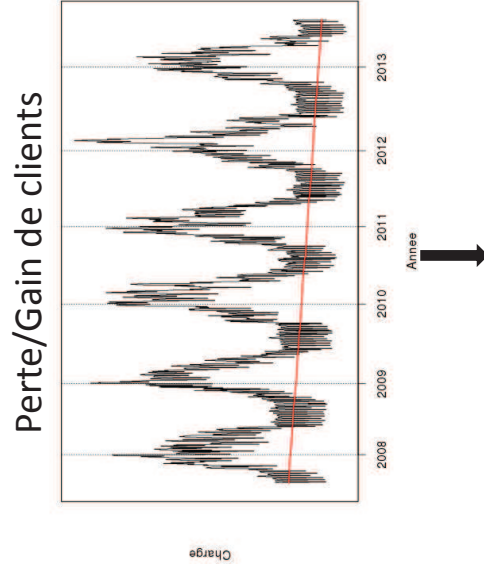
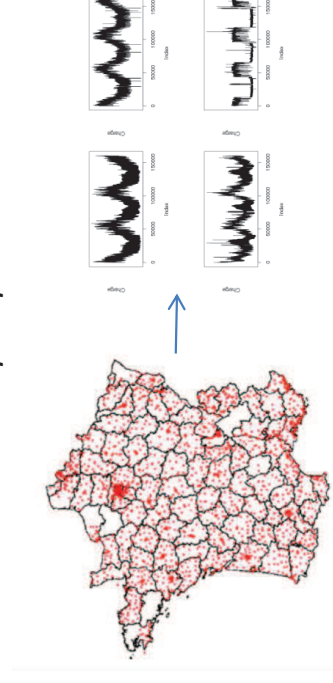
Consummation

Ouverture du marché de l'électricité



Développement de nouvelles technologies de mesure

Exemple: poste source



Besoin de méthodes adaptatives temporellement

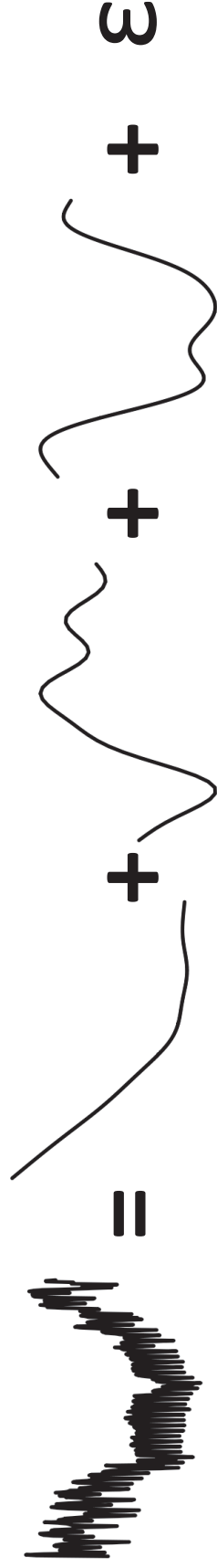
Besoin de développer des méthodes automatiques d'estimation de modèles statistiques

Quelle famille de modèles utiliser ?

- ▶ Modélisation non paramétrique $\Rightarrow$ *Beaucoup de variables explicatives*
- ▶ Modélisation paramétrique : modèle Météhore [3] $\Rightarrow$ *Beaucoup de paramètres à estimer et peu évolutif*
- ▶ Modèle additif (Hastie et Tibshirani [8]) :

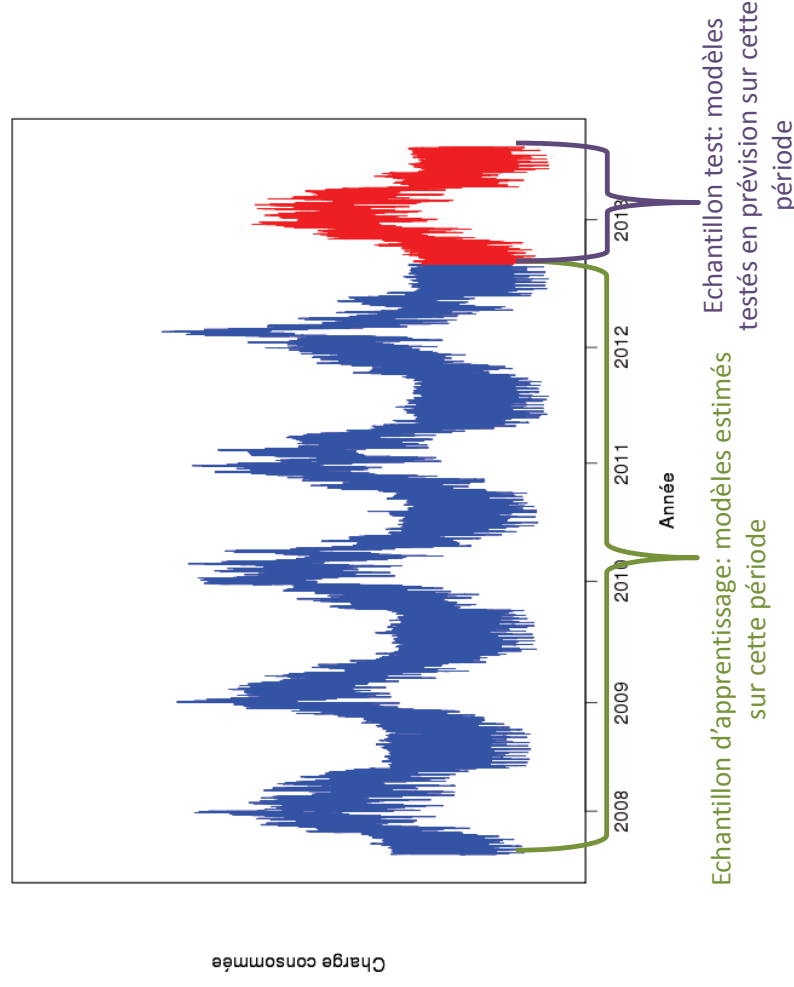
$$Y_i = \beta_0 + f_1(X_{1,i}) + \dots + f_d(X_{d,i}) + \dots + \varepsilon_i$$

- ▶ Application à la prévision de consommation électrique au niveau agrégé : Pierrot et Goude [13], Fan et Hyndman [6]
- ▶ Application à la prévision de consommation électrique au niveau local : Goude, Nedellec et Kong [7], Nedellec, Cugliari et Goude [12]



$$\text{Load} = f(T) + g(H) + h(D) + \varepsilon$$

Données du portefeuille d'un fournisseur



Modélisation de la consommation nationale

$$\begin{aligned}
Y_t = & s_1(Toy_t) + \sum_{i=1}^9 \sum_{j=1}^3 \alpha_{i,j} \mathbb{1}_{DayType_t=i} \mathbb{1}_{Offset_t=j} \\
& + \sum_{i=1}^{12} \beta_i T_{t-6} \mathbb{1}_{Month_t=i} + s_2(CC_t) + s_3(W_t) + s_4(T_t) \\
& + s_5(\theta_t^{0.996}) + s_6(\theta_t^{Max,0.983}) \\
& + s_7(\theta_t^{max,0.983}) + \gamma \Gamma_t + s_8(t) + \varepsilon_t
\end{aligned}$$

- ▶ T_{oy} : Moment dans l'année
- ▶ $DayType$: Type de jour
- ▶ $Offset$: Période homogène en termes de consommation
- ▶ T_{t-6} : Température retardée de 6 heures
- ▶ T : Température
- ▶ $\theta_t^{0.996}$: Lissage exponentiel de la température de coefficient 0.996
- ▶ Γ : Tendence associée à la part chauffage

Modélisation de la consommation nationale

$$\begin{aligned} Y_t = & s_1(Toy_t) + \sum_{i=1}^9 \sum_{j=1}^3 \alpha_{i,j} \mathbb{1}_{DayType_t=i} \mathbb{1}_{Offset_t=j} \\ & + \sum_{i=1}^{12} \beta_i T_{t-6} \mathbb{1}_{Month_t=i} + s_2(CC_t) + s_3(W_t) + s_4(T_t) \\ & + s_5(\theta_t^{0.996}) + s_6(\theta_t^{Max,0.983}) \\ & + s_7(\theta_t^{max,0.983}) + \gamma \Gamma_t + s_8(t) + \varepsilon_t \end{aligned}$$

- Modèle obtenu à l'aide de méthode stepwise à dire d'un expert
- Long, fastidieux, probablement non optimal et **difficilement généralisable à d'autres jeux de données**

Comment automatiser la sélection de variables et l'estimation des modèles ?

Modélisation de la consommation nationale

$$\begin{aligned} Y_t = & s_1(Toy_t) + \sum_{i=1}^9 \sum_{j=1}^3 \alpha_{i,j} \mathbb{1}_{DayType_t=i} \mathbb{1}_{Offset_t=j} \\ & + \sum_{i=1}^{12} \beta_i T_{t-6} \mathbb{1}_{Month_t=i} + s_2(CC_t) + s_3(W_t) + s_4(T_t) \\ & + s_5(\theta_t^{0.996}) + s_6(\theta_t^{Max,0.983}) \\ & + s_7(\theta_t^{max,0.983}) + \gamma \Gamma_t + s_8(t) + \varepsilon_t \end{aligned}$$

- Modèle obtenu à l'aide de méthode stepwise à dire d'un expert
- Long, fastidieux, probablement non optimal et **difficilement généralisable à d'autres jeux de données**

Comment automatiser la sélection de variables et l'estimation des modèles ?

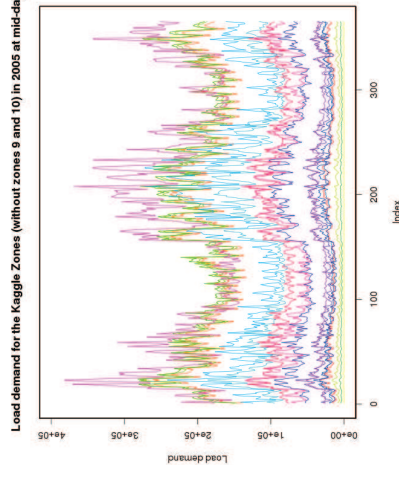
Etude d'une maille locale

Poste source



- ▶ Liaison entre les réseaux de transport et de distribution
- ▶ Environ 2200
- ▶ Forts enjeux industriels

Compétition Kaggle GEFCom 2012

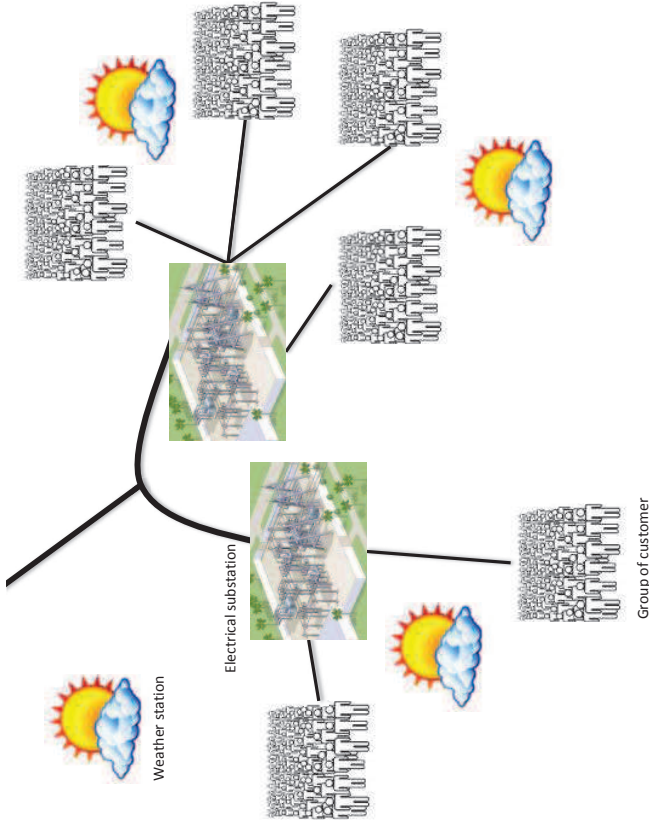


- ▶ Compétition présentée par Hong, Pinson et Fan [9]
- ▶ 17 zones électriques à prévoir, 11 stations météorologiques, aucune information géographique
- ▶ Version idéalisée mais riche en enseignements

Modélisation de la consommation locale

Modèle inspiré par Nedellec *et al.* [12] et par Goude *et al.* [7] :

$$Y_{t,j} = \sum_{q=1}^{11} m_q \mathbb{1}_{DayType_t=q} + h(T_{oy_t}) + k(t) + \sum_{i \in \mathbf{S}_j} l_i(T_{t,S_i}) + \sum_{i \in \mathbf{S}_j} g_i(\theta_{t,S_i}) + \varepsilon_t$$

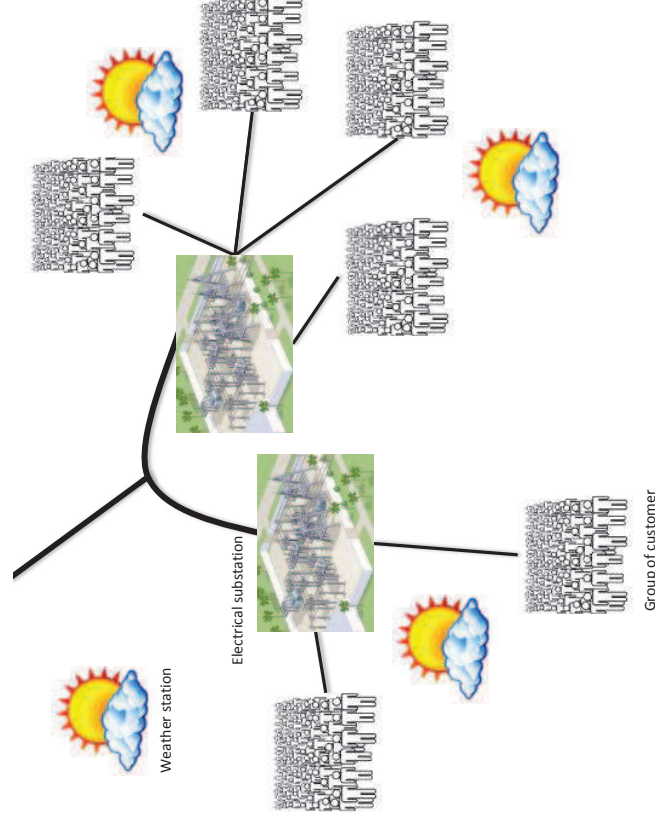


Comment procéder **automatiquement sans s'imposer un nombre fixe** de stations météorologiques ?

Modélisation de la consommation locale

Modèle inspiré par Nedellec *et al.* [12] et par Goude *et al.* [7] :

$$Y_{t,j} = \sum_{q=1}^{11} m_q \mathbb{1}_{\text{DayType}_t=q} + h(\text{To}_t) + k(t) + \sum_{i \in \mathbf{S}_j} l_i(T_{t,S_i}) + \sum_{i \in \mathbf{S}_j} g_i(\theta_{t,S_i}) + \varepsilon_t$$



Comment procéder **automatiquement sans s'imposer un nombre fixe** de stations météorologiques ?

Régression pénalisée

$$Y = \beta_0 + \sum_{k=1}^m \beta_{1,k} B_{1,k}^q(X_1) + \sum_{k=1}^m \beta_{2,k} B_{2,k}^q(X_2) + \sum_{k=1}^m \beta_{3,k} B_{3,k}^q(X_3) + \cdots + \sum_{k=1}^m \beta_{d,k} B_{d,k}^q(X_d) + \varepsilon$$

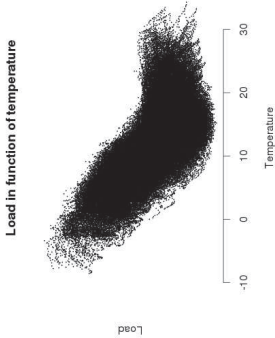
- Approximation : B-Splines (De Boor, 1978 [4])
- Estimation : P-Splines (Eilers et Marx, 1996 [5]) :
- Sélection : Group LASSO (Yuan et Lin, 2006 [14])

Régression pénalisée

$$Y = \beta_0 + \sum_{k=1}^m \beta_{1,k} B_{1,k}^q(X_1) + \sum_{k=1}^m \beta_{2,k} B_{2,k}^q(X_2) + \sum_{k=1}^m \beta_{3,k} B_{3,k}^q(X_3) + \dots + \sum_{k=1}^m \beta_{d,k} B_{d,k}^q(X_d) + \varepsilon$$

- Approximation : B-Splines (De Boor, 1978 [4])
- Estimation : P-Splines (Eilers et Marx, 1996 [5]) :

$$\tilde{\beta} = \arg \min_{\beta} \sum_{i=1}^n \left(Y_i - \beta_0 - \sum_{j=1}^d \sum_{k=1}^m \beta_{j,k} B_{j,k}^q(X_{i,j}) \right)^2 + \sum_{j=1}^d \lambda_j \|D_k \beta_j\|_2^2$$

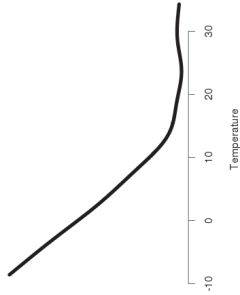


Effet lissé



Effect of temperature

Effect of temperature on load



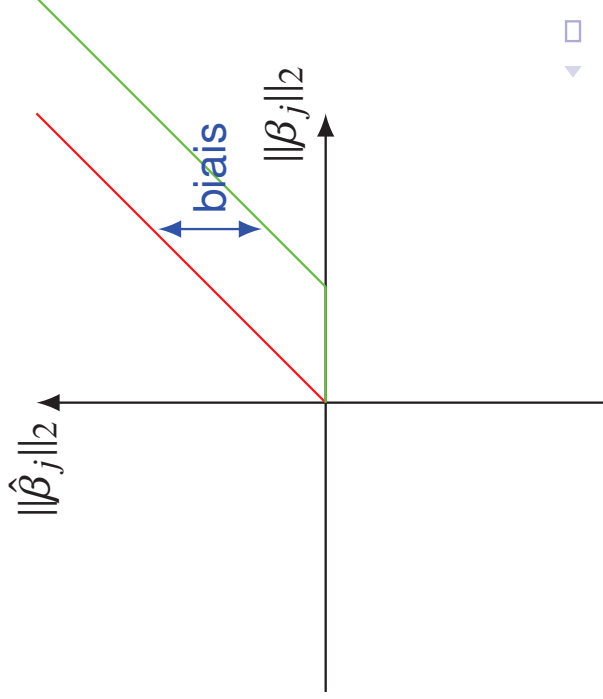
Estimateur P-Splines

Régression pénalisée

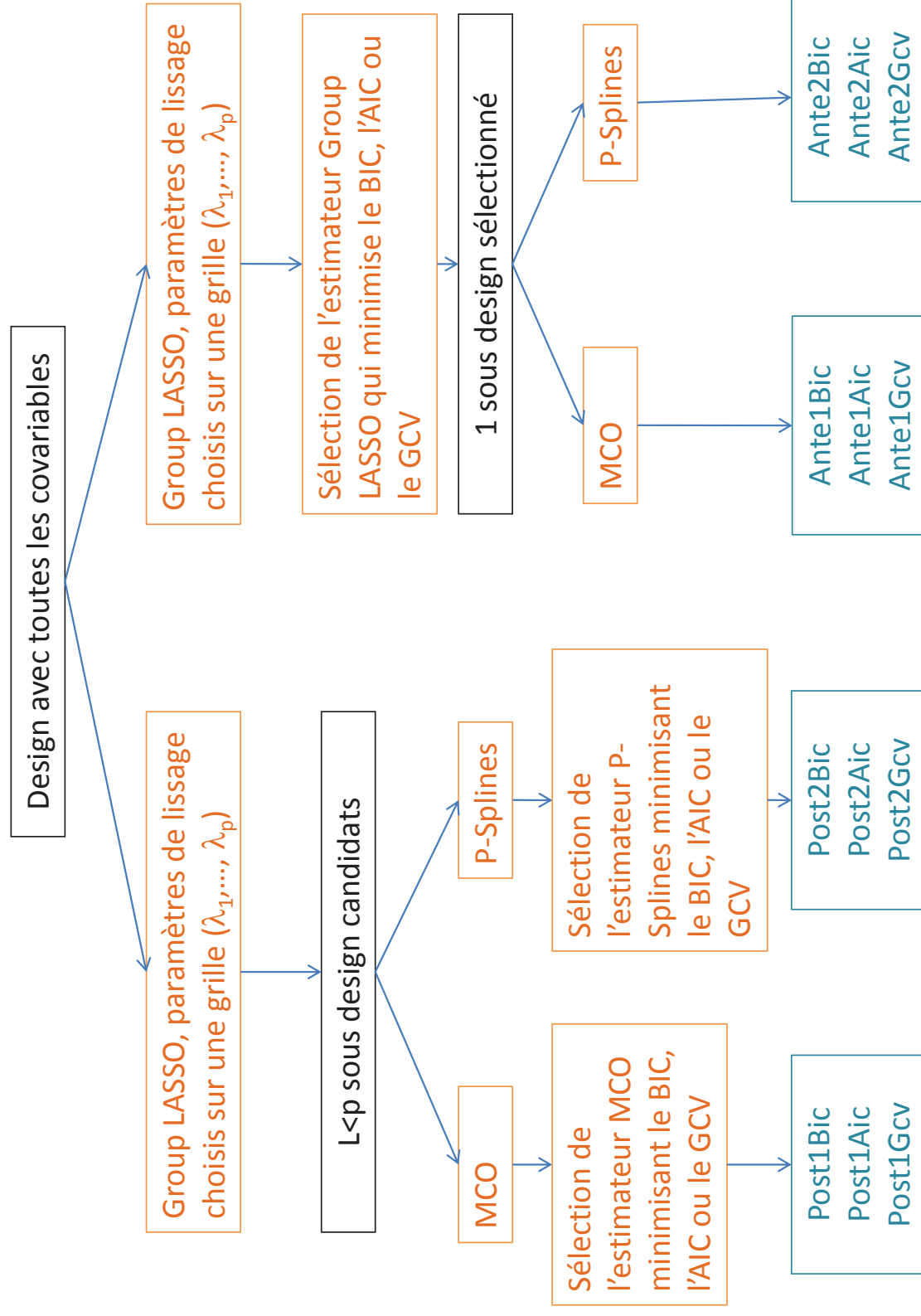
$$Y = \beta_0 + \sum_{k=1}^m \beta_{1,k} B_{1,k}^q(X_1) + \sum_{k=1}^m \beta_{2,k} B_{2,k}^q(X_2) + \sum_{k=1}^m \beta_{3,k} B_{3,k}^q(X_3) + \dots + \sum_{k=1}^m \beta_{d,k} B_{d,k}^q(X_d) + \varepsilon$$

- Approximation : B-Splines (De Boor, 1978 [4])
- Estimation : P-Splines (Eilers et Marx, 1996 [5]) :
- Sélection : Group LASSO (Yuan et Lin, 2006 [14]) :

$$\hat{\beta} = \arg \min_{\beta} \sum_{i=1}^n \left(Y_i - \beta_0 - \sum_{j=1}^d \sum_{k=1}^m \beta_{j,k} B_{j,k}^q(X_{i,j}) \right)^2 + \lambda \sqrt{m} \sum_{j=1}^d \|\beta_j\|_2$$



Estimateurs en plusieurs étapes



Résultat asymptotique pour Post1Bic



Théorème

Sous nos hypothèses, soit $\lambda_n \sim \sqrt{n \log(m_n d) / m_n}$, alors

$$P(S_{\lambda_n} = S^*) \rightarrow 1,$$

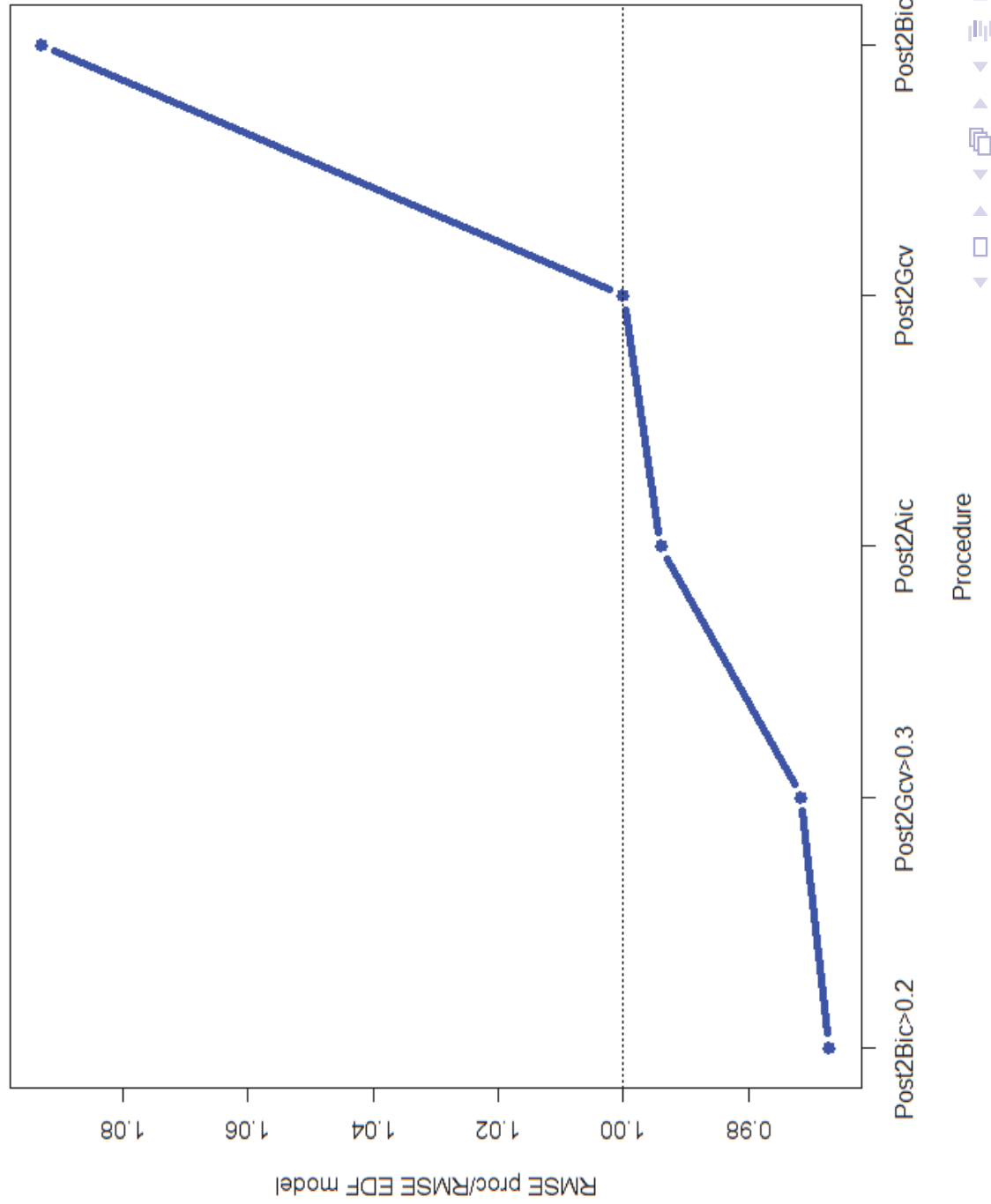
et pour tout λ tel que $S_\lambda \neq S^*$ et $\text{card}(S_\lambda) \leq M$

$$P(BIC(\lambda) - BIC(\lambda_n) > 0) \rightarrow 1.$$

Voir Antoniadis, Goude, Poggi et Thouvenot [1] pour la démonstration. De plus, sous les hypothèses posées [1], l'estimateur **converge en erreur quadratique moyenne**.

Performance à moyen terme

Ratio between the RMSE of the procedures and EDF model



Application locale - Détails de l'application GEFCom 2012

- ▶ Modèles **moyen terme** appris en 2004-2007, testés sur les 6 premiers mois de 2008
- ▶ Modèles horaires inspirés de Nedellec *et al.* [12] :

$$Y_{t,j} = \sum_{q=1}^{11} m_q \mathbb{1}_{\text{DayType}_t=q} + h(T_{\text{oy}_t}) + k(t) + \sum_{i \in \mathbf{S}_i^{\mathbf{j}}} l_i(T_{t,S_i}) + \sum_{i \in \mathbf{S}_i^{\mathbf{j}}} g_i(\theta_{t,S_i}) + \varepsilon_t$$

- ▶ A chaque station météorologique, des températures brutes et lissées sont associées
- ▶ Quelle stratégie adopter face au nombre de stations météorologiques ?
- ▶ Application de Post2Bic et Post2Aic

Passage aux 17 zones

Trois modèles de référence :

Modèle saturé

$$Y_{t,j} = \sum_{q=1}^{11} m_q \mathbb{1}_{DayType_t=q} + h(T_{Oy_t}) + k(t) + \sum_{b=1}^{11} l_b(T_{t,b}) + \sum_{b=1}^{11} g_b(\theta_{t,b}) + \varepsilon_t$$

Nombre fixe de stations météorologiques sélectionnées, s’inspire de Lloyd [11]

Modèle avec un signal de température

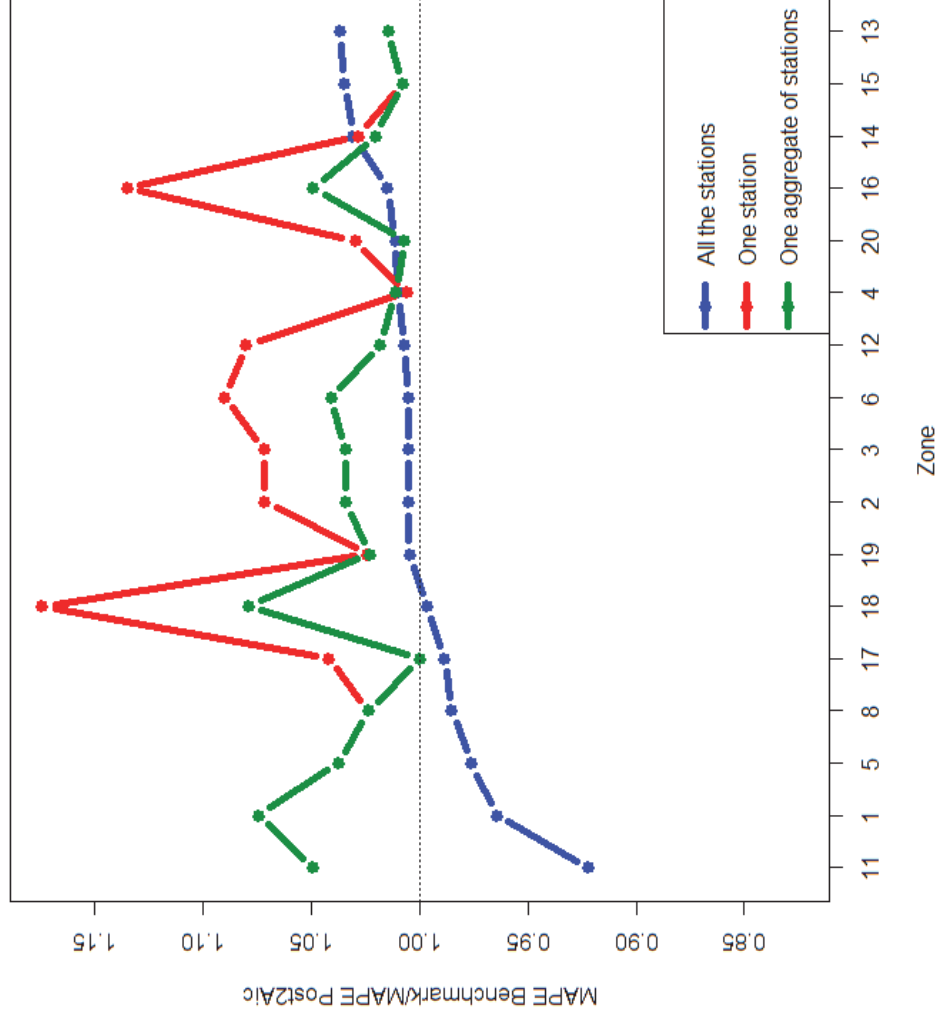
$$Y_{t,j} = \sum_{q=1}^{11} m_q \mathbb{1}_{DayType_t=q} + h(T_{Oy_t}) + k(t) + l(T_{t,j}) + g(\theta_{t,j}) + \varepsilon_t$$

- ▶ Sélection d’une station météorologique minimisant le MAPE sur un échantillon de validation : **nombre fixe** de stations météorologiques, s’inspire de Ben Taieb *et al.* [2]
- ▶ Sélection de type pas à pas de moyennes agrégées de températures : **nombre variable** de stations météorologiques, s’inspire de Hong *et al.* [10]





Ratio between the MAPE of Benchmarks and the MAPE of Post2Aic



- ▶ Représentation des MAPE des modèles de référence relativement à celui de Post2Aic
- ▶ Valeur 1 $\Rightarrow$ le MAPE du modèle considéré = MAPE de Post2Aic
- ▶ Valeur $< 1 \Rightarrow$ le modèle de référence a des meilleures performances prédictives d'après le critère MAPE

Apport des Mathématiques

- ▶ Modélisation et son approximation
- ▶ Algorithmes d'optimisation pour estimer les coefficients associés à la modélisation
- ▶ Obtention de méthodes ayant automatiquement des performances équivalentes à des méthodes manuelles
- ▶ Assurances théoriques sur la justesse en sélection et en estimation des estimateurs obtenus
- ▶ Mesure de la qualité des prévisions

Sommaire

Introduction

Prévision de consommation électrique

Protection de données personnelles

Des données riches mais à protéger

Vers une définition de l'anonymisation

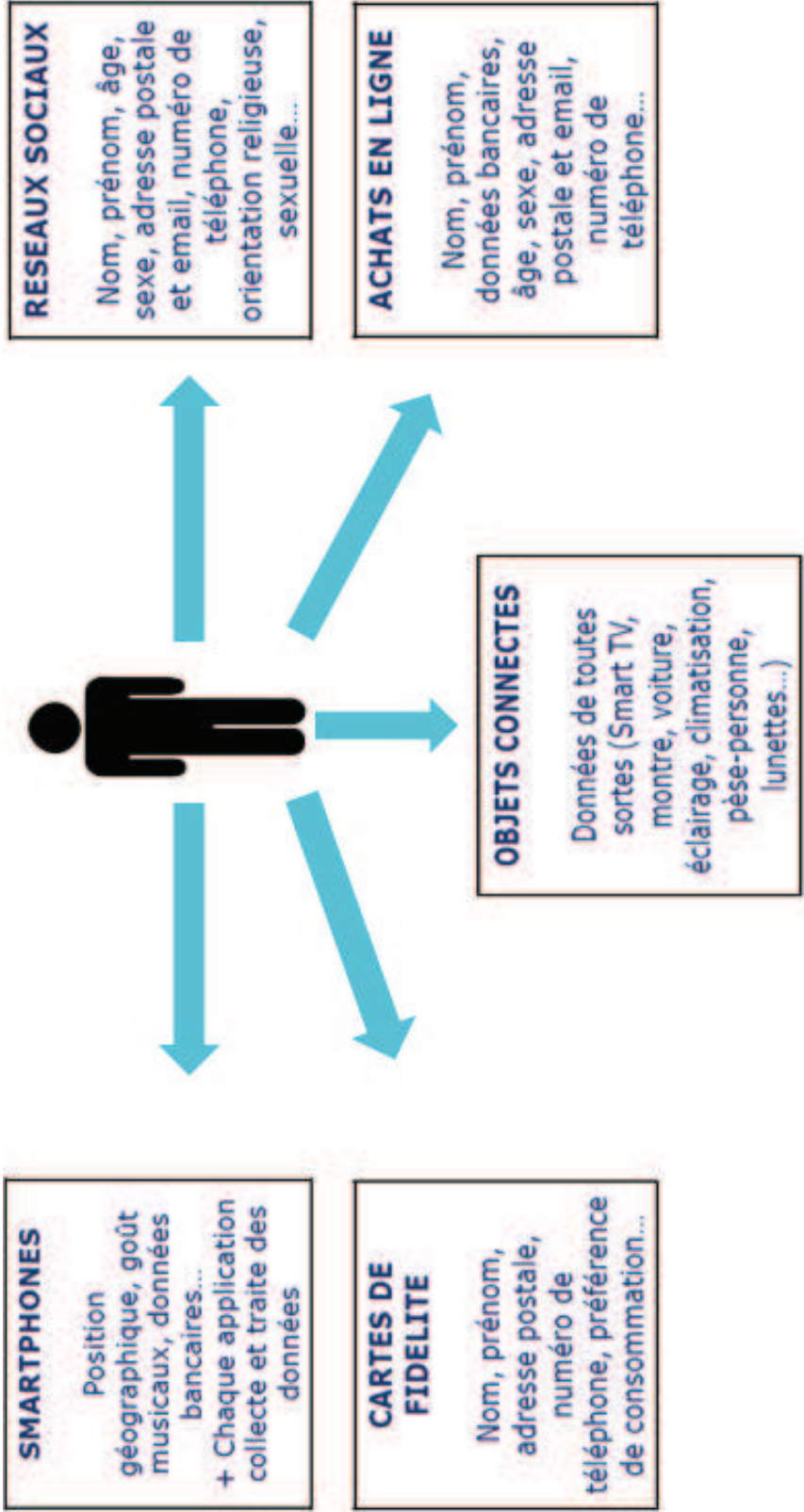
Techniques de dé-identification et de ré-identification

Exemple

Un monde de données

Conclusion

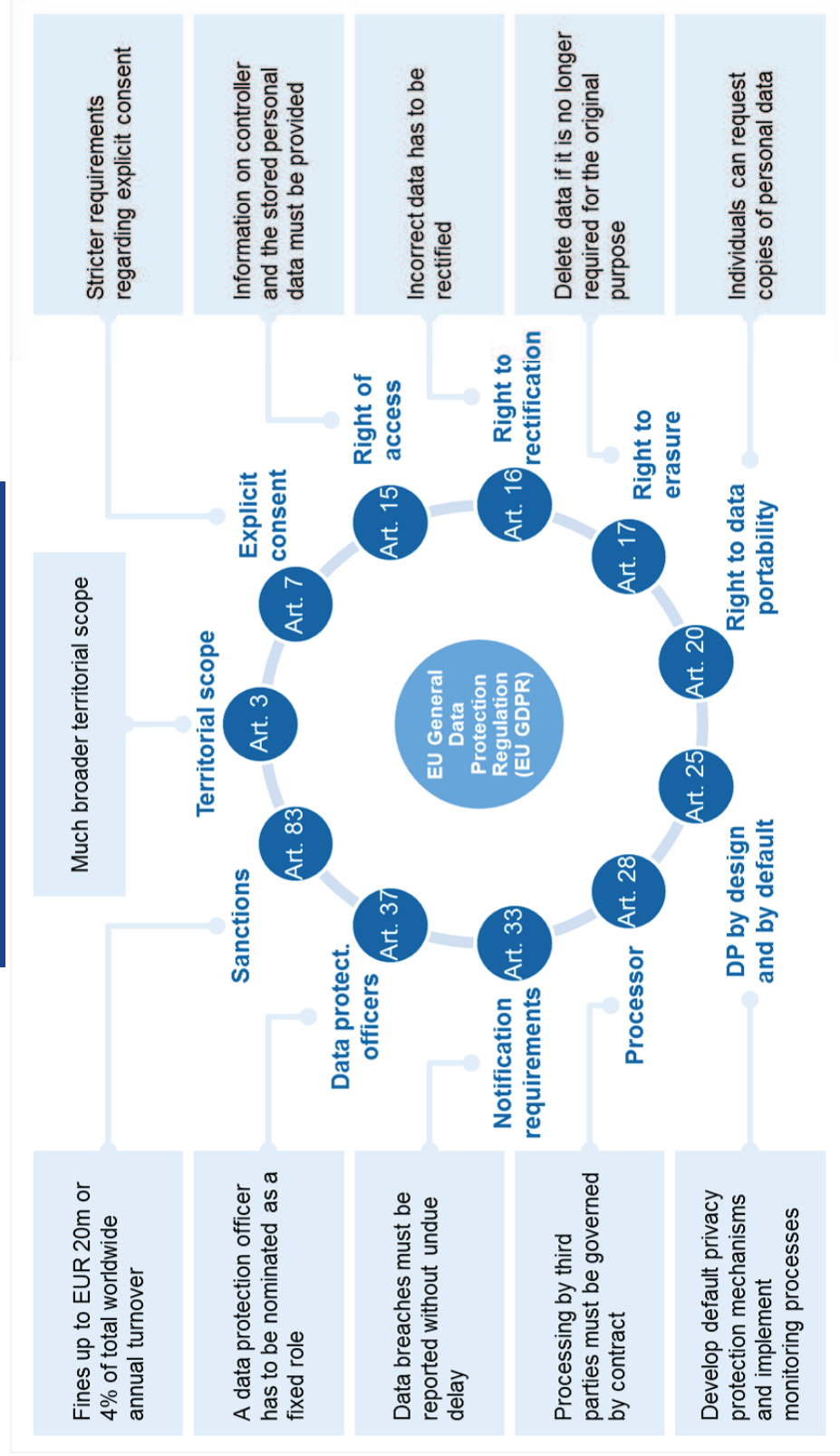
Données personnelles



Les données, une mine d'or

- ▶ Transformation numérique/digitale
 - ▶ Désilotage des données
 - ▶ Nouveaux marchés
- ▶ Open Data
 - ▶ Forte exigence réglementaire
 - ▶ Favorise l'ouverture et la circulation des données et du savoir
 - ▶ Permet de garantir un environnement numérique ouvert
 - ▶ Doit être respectueux de la vie privée
- ▶ Big Data
 - ▶ Permet de nouvelles méthodes de croisement de données et donc de ré-identifier des individus

Cadre réglementaire - RGPD



Protection des données personnelles ?

Protéger les données personnelles, c'est ...

- ▶ Empêcher un individu d'être isolé
- ▶ Empêcher la possibilité de corréler les informations sur un individu de plusieurs jeux de données
- ▶ Empêcher d'obtenir des informations sur les individus grâce à des variables exogènes

...et concerne

- ▶ des informations identifiant directement ou quasi directement des individus (nom, numéro de téléphone, adresse, etc.)
- ▶ des informations sensibles portant sur les opinions politiques ou religieuses, la santé et les habitudes, etc. des individus

Protéger les données par la pseudonymisation

La pseudonymisation

- ▶ consiste à chiffrer ou supprimer les éléments directement identifiants...
- ▶ ...et est souvent **insuffisante** car des éléments peuvent être croisées ou des éléments extérieures pour permettre la ré-identification
- ▶ Par exemple AOL en 2006
 - ▶ Mise en Open Data de 20 millions de requêtes de recherche dont l'identifiant des utilisateurs avaient été chiffrés
 - ▶ Par recoupement géographique et utilisation d'informations se retrouvant avec les recherches effectuées (centre d'intérêt, âge, etc.), certains utilisateurs ont été ré-identifiés
 - ▶ Le CTO et plusieurs salariés d'AOL renvoyés suite à cette mise en Open Data

Anonymisation : Réalisation d'un compromis utilité/protection

Anonymisation

=



Protection de la vie privée

Conservation de l'utilité des données

L'anonymisation est fortement dépendante

- ▶ De la typologie des données $\Rightarrow$ données d'une table médicale $\neq$ données de réseaux sociaux par exemple
- ▶ Des traitements futures $\Rightarrow$ besoin de données fines individuellement ou seulement de groupes d'individus ?
- ▶ Du temps $\Rightarrow$ de nouveaux jeux de données et des nouvelles méthodes permettent de ré-identifier

L'anonymisation, un sujet complexe

- ▶ Pas de solution universelle
 - ▶ Réalisation d'un compromis utilité/protection
 - ▶ Fortement dépendante du **type de données**, du **type d'application** et du **temps**
- ▶ Pas de risque zéro avec l'anonymisation
 - ▶ Implique l'étude du risque ré-identification
 - ▶ Risque résiduelle doit être faible par rapport aux bénéfices
- ▶ Procédé global
 - ▶ Comprenant des composants de **dé-identification**, de **ré-identification** et d'**évaluation de la qualité des données**
 - ▶ Auditable

Généralité sur la dé-identification - Trois familles de techniques

- ▶ L'anonymisation peut se partager en trois familles de techniques
 - ▶ Transformer les données pour avoir des individus fictifs
 - ▶ Agréger et généraliser les individus pour fournir des agrégats symbolisant des groupes d'individus
 - ▶ Générer entièrement des données fictives
- ▶ Pour plus de protection, les techniques peuvent être combinées

Généralité sur la dé-identification - Techniques de dé-identification

- ▶ Techniques de permutation et de puzzling
 - ▶ Approches déstructurant et modifiant l'architecture des données
 - ▶ Individus ainsi créés peuvent être irréalistes
- ▶ Application de bruits additifs ou multiplicatifs
 - ▶ Ajout de bruit : ré-identification par filtrage, MAP, etc.
 - ▶ Bruits multiplicatifs : ré-identification possible lorsque l'attaquant connaît un sous-échantillon
 - ▶ Confidentialité différentielle
- ▶ Généralisation des similaires
- ▶ Microagrégation

Généralité sur la dé-identification - Techniques de dé-identification

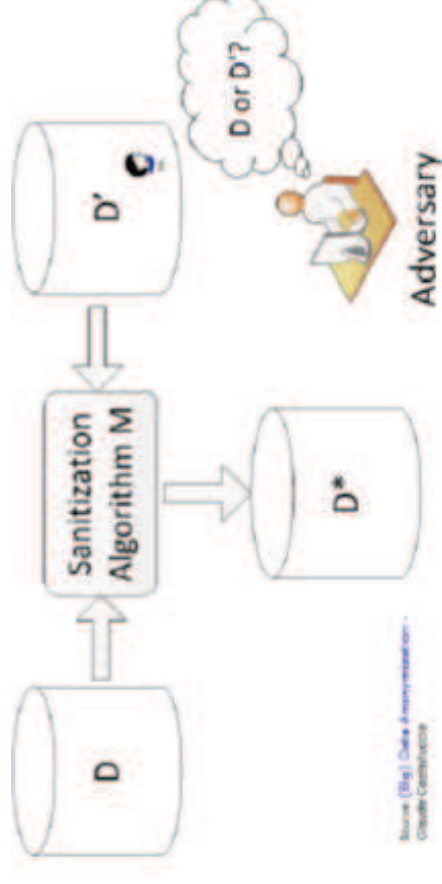
- ▶ Techniques de permutation et de puzzling
 - ▶ Approches déstructurant et modifiant l'architecture des données
 - ▶ Individus ainsi créés peuvent être irréalistes
- ▶ Application de bruits additifs ou multiplicatifs
 - ▶ Ajout de bruit : ré-identification par filtrage, MAP, etc.
 - ▶ Bruits multiplicatifs : ré-identification possible lorsque l'attaquant connaît un sous-échantillon
 - ▶ **Confidentialité différentielle**
- ▶ **Généralisation des similaires**
- ▶ **Microagrégation**

Confidentialité différentielle - Principe [?]

Notons ε un réel positif, M un mécanisme d'anonymisation, Im_M l'image de M . Alors M est dit ε -différentiellement confidentielle si pour tout D_1 et D_2 jeux de données adjacents et pour tout $S \subset Im_M$:

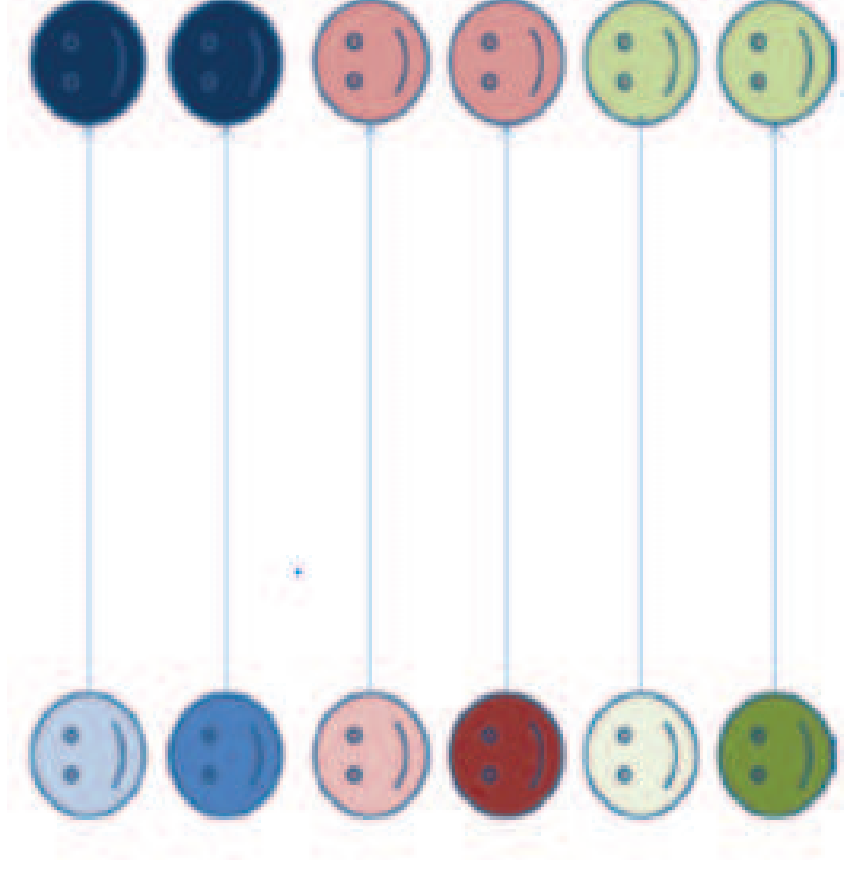
$$P\{M(D_1) \in S\} \leq \exp(\varepsilon)P\{M(D_2) \in S\}$$

- ▶ La définition assure que la présence ou l'absence d'un individu n'influencera pas significativement la sortie du mécanisme
- ▶ Plus ε est petit, plus il est compliqué de distinguer la présence ou l'absence d'un individu
- ▶ Lorsqu'il est nul, appliqué aux deux jeux de sortie adjacents, le mécanisme a la même probabilité de produire la sortie S . On ne peut donc plus distinguer la présence ou l'absence de l'individu différent



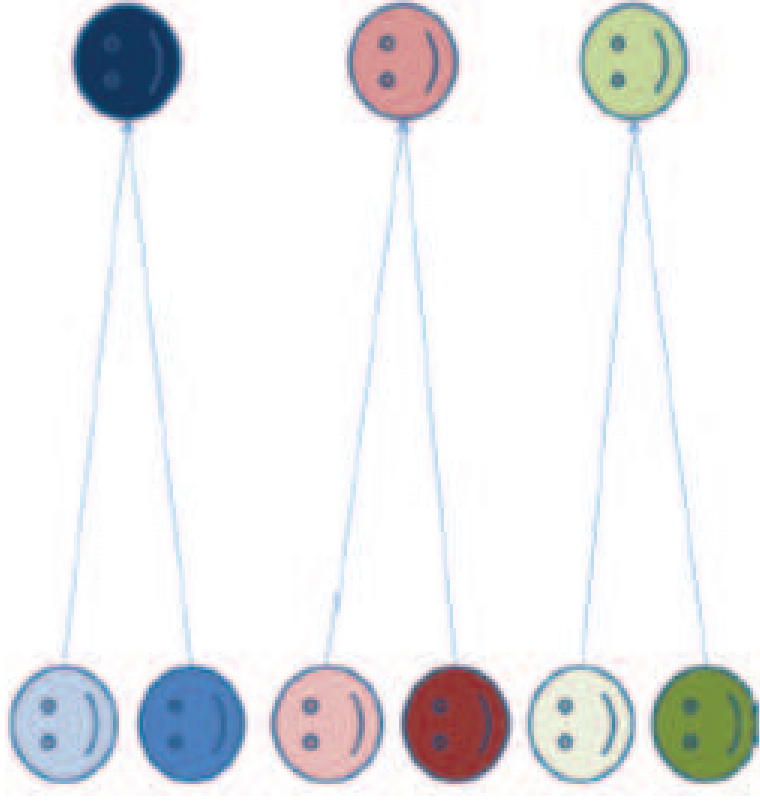
Généralisation des similaires - Principe

- ▶ K-anonymisation [?] : un jeu de données est K-anonyme si au moins K enregistrements existent pour chaque combinaison de quasi-identifiant
- ▶ L-Diversité [?] : un jeu de données respecte la L-Diversité si au moins K enregistrements existent pour chaque combinaison de quasi-identifiant et si dans chaque classe, la variable sensible prend au moins L valeurs différentes
- ▶ T-Proximité [?] : un jeu de données respecte la L-Diversité si au moins K enregistrements existent pour chaque combinaison de quasi-identifiant et si la répartition des valeurs sensibles respectent la répartition de celles-ci dans la population totale



Micro-agrégation - Principe

- ▶ Objectif : traiter différents types de variables (continues, catégorielle, etc.) et agréger les individus de manière la plus optimale possible.
- ▶ Deux phases
 - ▶ Partitionnement des données en des sous espaces de taille k au minimum
 - ▶ Application d'un opérateur d'agrégation
- ▶ Paramètres
 - ▶ Le choix de la taille k minimum
 - ▶ L'algorithme et les distances utilisés pour partitionner les données
 - ▶ Le choix de la méthode d'agrégation
- ▶ Risques de ré-identification : déduction d'informations probables sur un individu
- ▶ Utilité des données : métriques d'homogénéité (indices de Bouldin-Davies, Silhouette, etc.) et indicateurs de qualité globaux (RMSE, MAPE, etc.)



Ré-identification

- ▶ Plusieurs scénarios de ré-identification sont possibles
 - ▶ Isoler et identifier un ou quelques individus
 - ▶ Identifier l'ensemble des individus
 - ▶ Dédurre des informations sensibles sur un ou plusieurs individus
- ▶ Plusieurs catégories de méthodes de ré-identification
 - ▶ Utilisation d'un jeu de données non anonymisé pour ré-identifier un jeu de données anonymisé
 - ▶ Par exemple, utilisation d'un réseau social non anonymisé pour attaquer un réseau social anonymisé
 - ▶ Utilisation d'informations présentes dans le jeu de données anonymisée
 - ▶ Par exemple, le contenu pour attaquer un réseau social anonymisé
- ▶ Utilisation de méthodes classiques de Machine Learning et d'optimisation combinatoire

Apport des Mathématiques

- ▶ Cryptologie et chiffrement
- ▶ Pré-traitement des données
- ▶ Etude de Probabilité
- ▶ Algorithmes de partitionnement et de généralisation
- ▶ Métriques de perte d'information
- ▶ Mesure de la “Privacy”

Sommaire

Introduction

Prévision de consommation électrique

Protection de données personnelles

Un monde de données

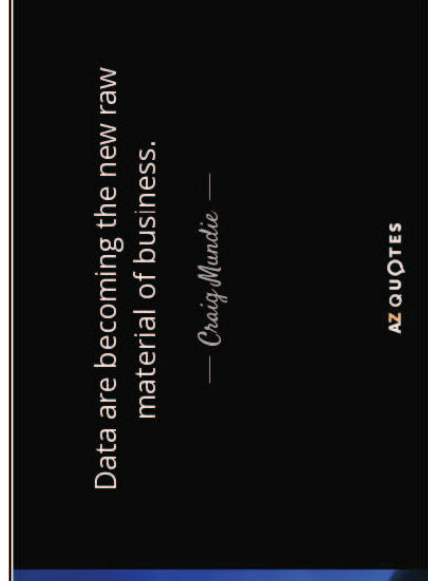
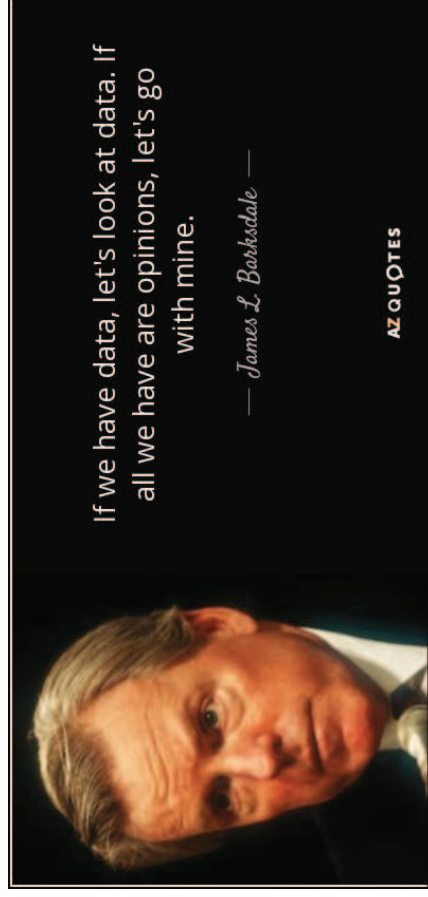
Citations

Big Data

Apport des Mathématiques

Conclusion

Des données au Machine Learning



Des données au Machine Learning



Data isn't information; information isn't knowledge; knowledge isn't wisdom.

— Jan Leve —

AZ QUOTES



The goal is to turn data into information, and information into insight.

— Carly Fiorina —

AZ QUOTES



We used to be calorie poor and now the problem is obesity. We used to be data poor, now the problem is data obesity.

— Hal Varian —

AZ QUOTES

V2SOFT

“Hiding within those mounds of data is knowledge that could change the life of a patient, or change the world”

- Atul Butte,
(Stanford School of Medicine)

Des données au Machine Learning



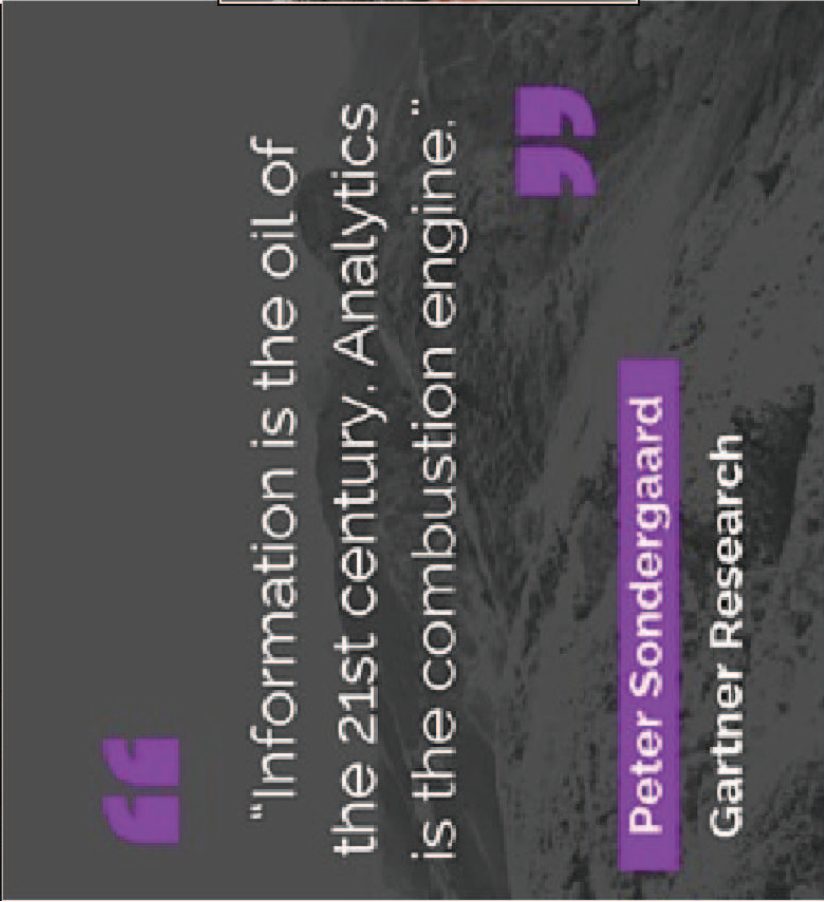
If you torture the data long enough,
it will confess.

— *Ronald Coase* —

AZ QUOTES

**“Data is the sword of
the 21st century,
those who wield it
well, the Samurai.”**

— Jonathan Rosenberg
SVP, Product Management
Google



Peter Sondergaard
Gartner Research



A breakthrough in machine learning
would be worth ten Microsofts.

— *Bill Gates* —

AZ QUOTES

Des données au Machine Learning



— Thomas H. Davenport —

AZ QUOTES



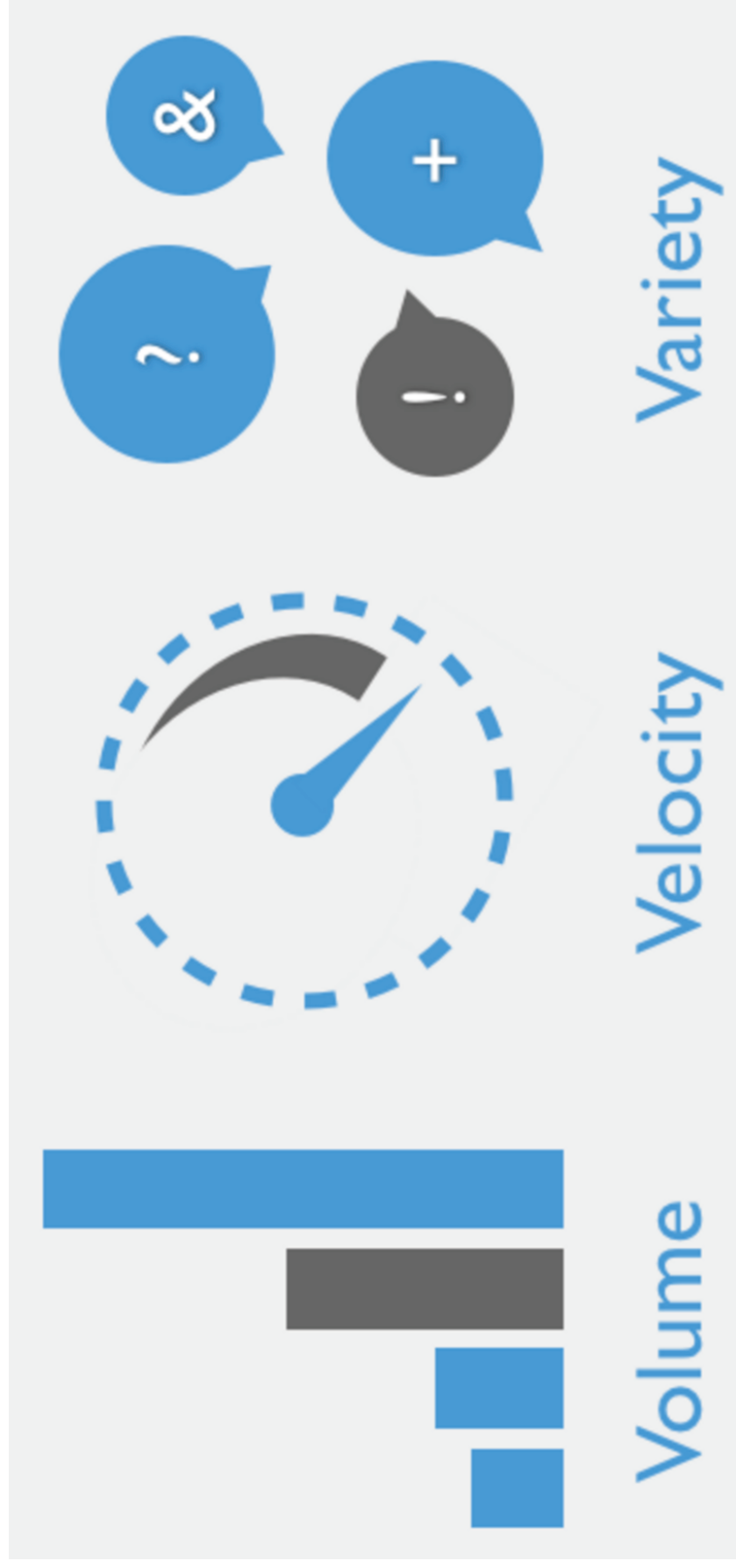
— Chris Lynch —

AZ QUOTES

Le Big Data : Pourquoi ?



Le Big Data : les 3 V



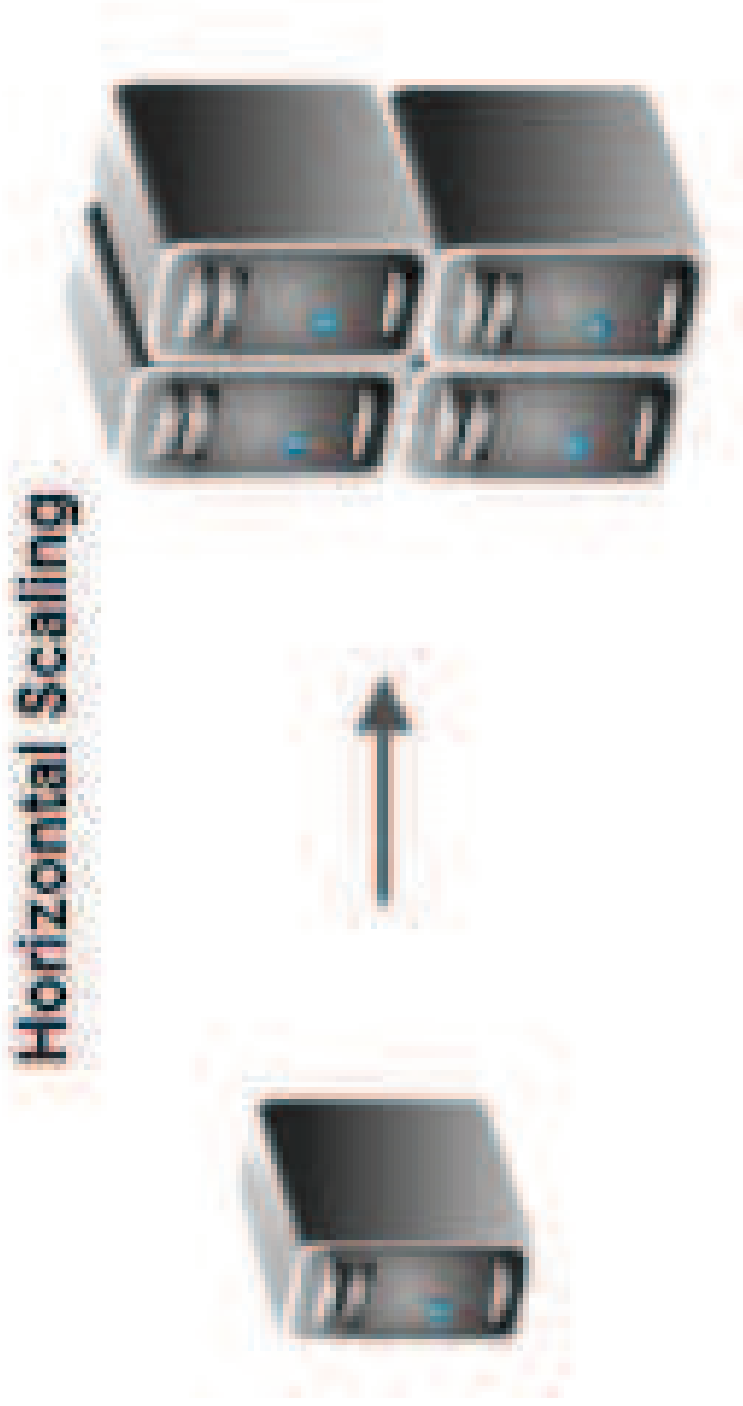
Le Big Data n'est pas ...



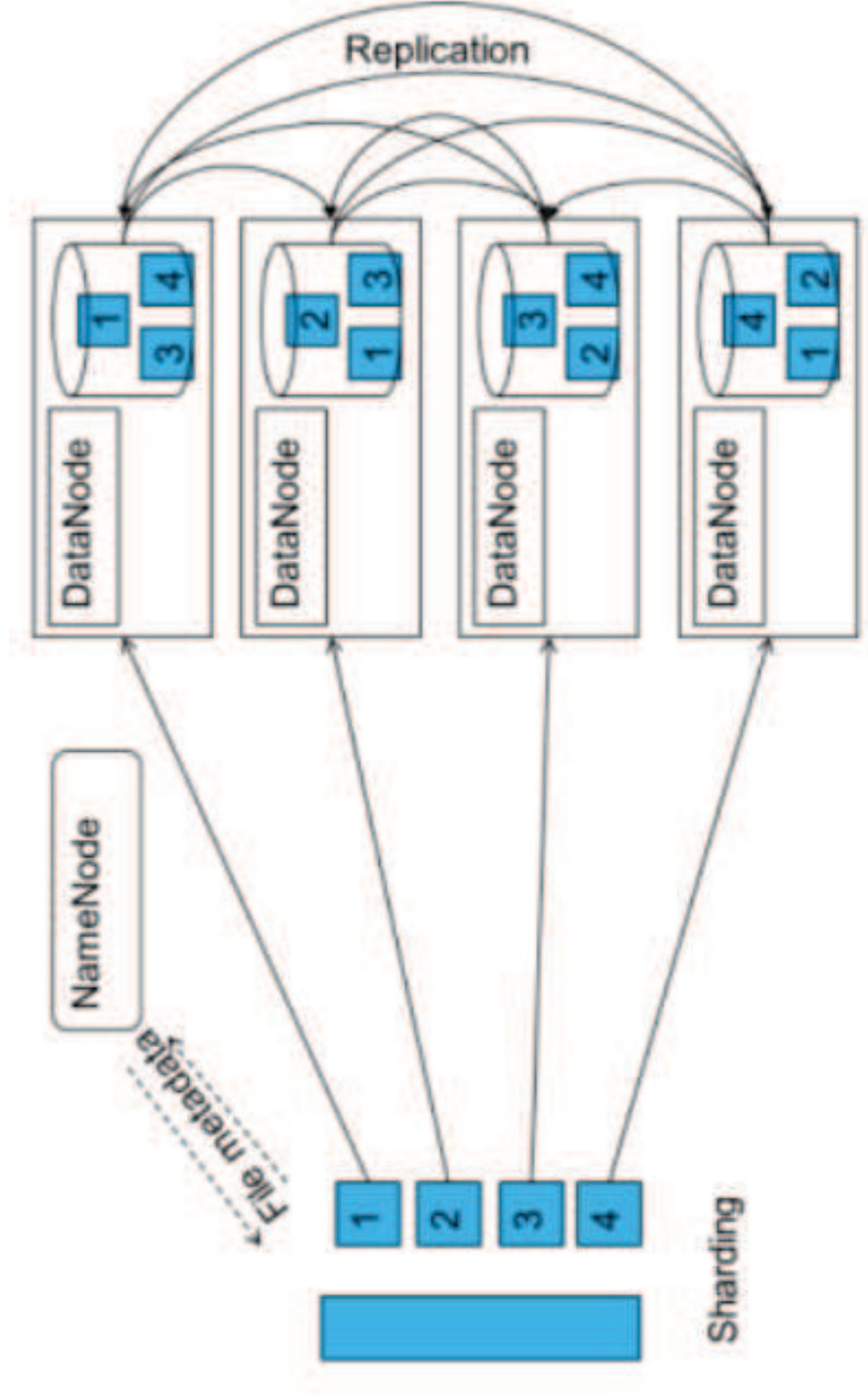
Utilisation de technologies Big Data : scaling horizontal

Lorsque le volume des données augmente, augmentation du nombre de serveurs $\Rightarrow$ plus souple et moins coûteux qu'un hyper calculateur

Horizontal Scaling

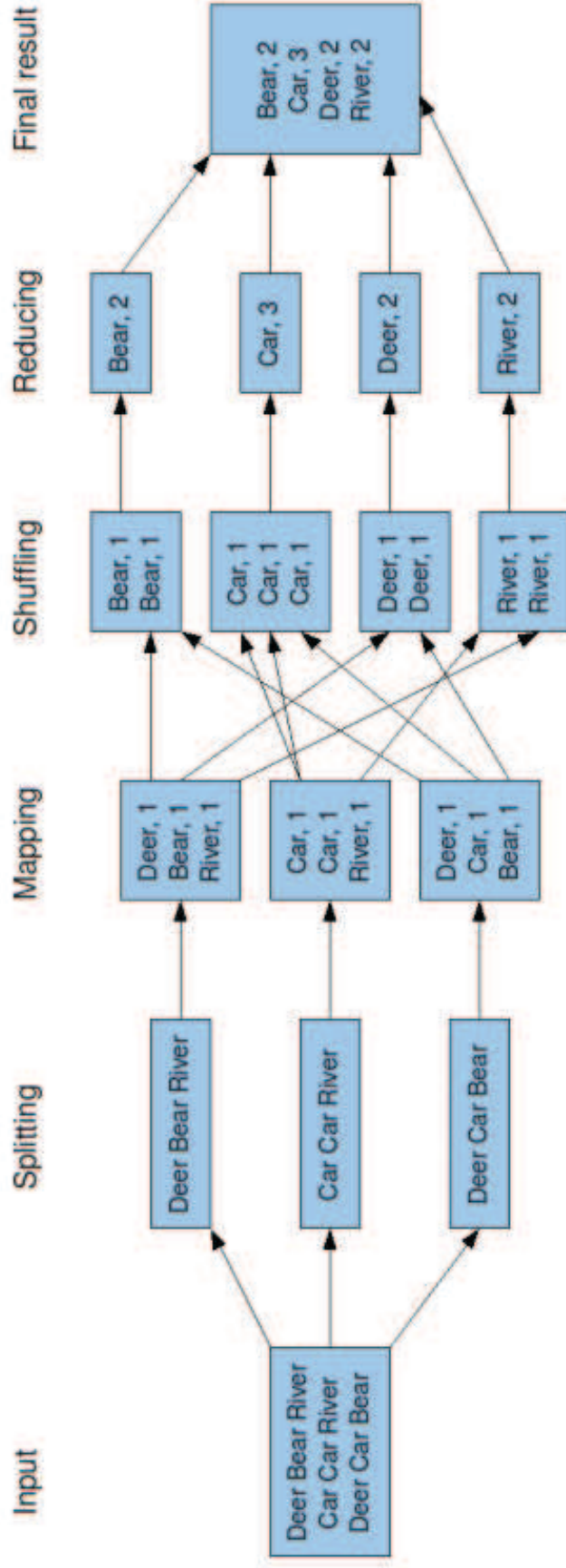


Système de fichiers distribués



Map Reduce

The overall MapReduce word count process



Apport des Mathématiques

- ▶ Le Big Data met en jeu un grand nombre de domaines des mathématiques appliquées : statistiques au sens large, traitement du signal et des images, optimisation convexe et non convexe, probabilités appliquées (graphes aléatoires, algorithmes stochastiques...).
- ▶ Toutes les techniques statistiques doivent être passées à l'échelle. Par exemple, si on considère la recherche de la solution de l'estimation des MCO d'une régression linéaire sous hypothèses classiques
 - ▶ Classiquement un produit matricielle impliquant une inversion de matrice $\Rightarrow$ Comment l'obtenir dans un contexte de données distribuées ?
 - ▶ Utilisation d'un algorithme de descente de gradient stochastique par exemple en contexte Big Data

Sommaire

Introduction

Prévision de consommation électrique

Protection de données personnelles

Un monde de données

Conclusion

- ▶ Les Mathématiques appliquées aux données, utilisant notamment des Statistiques et des Probabilités appliquées, des algorithmes d'optimisation, de recommandation, de traitement du signal, etc., sont de plus en plus demandées
 - ▶ Dans tous les secteurs : banque et assurance, industrie lourde et pharmaceutique, Transport, Energie, etc.
 - ▶ Pour de nombreuses applications : cyber sécurité, connaissance client (comportement et ciblage), maintenance, dimensionnement de réseau, prévision de série temporelle, gestion de risque, etc.
- ▶ Nombreuses sources de recrutement : ENSAE, ENSAI, ISUP, Université, Masters spécialisés d'école d'ingénieur, etc.

- [1] A. ANTONIADIS, Y. GOUDE, J.-M. POGGI, AND V. THOUVENOT, *Sélection de variables dans les modèles additifs avec des estimateurs en plusieurs étapes*, tech. report. <https://hal.archives-ouvertes.fr/hal-011116100>,.
- [2] S. BEN TAIEB AND R. HYNDMAN, *A gradient boosting approach to the kaggle load forecasting competition*, International Journal of Forecasting, 30 (2014), pp. 382 – 394.
- [3] 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).
- [4] C. DE BOOR, *A practical guide to splines*, Springer-Verlag New York, 1978.
- [5] P. EILERS AND B. MARX, *Flexible smoothing with b-splines and penalties*, Statistical Science, 11 (1996), pp. 89–121.
- [6] S. FAN AND R. HYNDMAN, *Short-term load forecasting based on a semi-parametric additive model*, Power Systems, IEEE Transactions on, 27 (2012), pp. 134–141.
- [7] 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 (2014), pp. 440–446.
- [8] T. J. HASTIE AND R. J. TIBSHIRANI, *Generalized additive models*, London : Chapman & Hall, 1990.
- [9] T. HONG, P. PINSON, AND S. FAN, *Global energy forecasting competition 2012*, International Journal of Forecasting, 30 (2014), pp. 357 – 363.
- [10] T. HONG, P. WANG, AND L. WHITE, *Weather station selection for electric load forecasting*, International Journal of Forecasting, (To appear).
- [11] J. LLOYD, *Gefcom2012 hierarchical load forecasting : Gradient boosting machines and gaussian processes*, International Journal of Forecasting, 30 (2014), pp. 369–374.
- [12] R. NEDELLEC, J. CUGLIARI, AND Y. GOUDE, *Gefcom2012 : Electric load forecasting and backcasting with semi-parametric models*, International Journal of Forecasting, 30 (2014), pp. 375 – 381.
- [13] A. PIERROT AND Y. GOUDE, *Short-term electricity load forecasting with generalized additive models*, Proceedings of ISAP power, (2011), pp. 593–600.
- [14] M. YUAN AND Y. LIN, *Model selection and estimation in regression with grouped variables*, Journal of the Royal Statistical Society, Series B, 68 (2006), pp. 49–67.

