

Table des matières

Introduction	7
1 Contexte astrophysique	13
1.1 Brève histoire de la classification spectrale	14
1.2 Classification spectrale et paramètres stellaires fondamentaux	18
1.3 À propos des données astrophysiques étudiées	20
2 Méthodes de traitement statistique de données	25
2.1 Analyse en Composantes Principales	26
2.1.1 Principe	27
2.1.2 Mise en œuvre	29
2.1.3 Propriétés	30
2.2 Regression inverse par tranches (SIR)	31
2.2.1 Principe	32
2.2.2 Mise en œuvre	34
2.2.3 Propriétés	36
2.2.4 Conditionnement de la matrice de covariance	36
2.3 Moindres carrés partiels	37
2.3.1 Principe	37
2.3.2 Mise en œuvre	38
2.3.3 Propriétés	39

2.4	MATISSE	39
2.4.1	Principe	40
2.4.2	Mise en œuvre	40
2.4.3	Propriétés	41
2.5	Conclusion	41
3	Régression basée sur l'ACP	43
3.1	Description de l'utilisation de l'ACP pour des problèmes de régression . . .	43
3.1.1	Choix du nombre de composantes	43
3.1.2	Choix de la méthode d'estimation	46
3.2	Particularité du traitement de données non-linéaires	48
3.2.1	Régression non-linéaire	48
3.2.2	Régression linéaire locale	49
3.2.3	Estimation basée sur les k-plus proches voisins	51
3.2.4	Sélection de directions pertinentes	51
3.3	Validation des approches	53
3.3.1	Cas d'un problème linéaire	54
3.3.2	Cas d'un problème non-linéaire	60
3.3.3	Cas d'un problème complexe	66
3.4	Application aux données astrophysiques	67
3.5	Sélection de directions sur données observées	70
3.5.1	Température effective (T_{eff})	70
3.5.2	Gravité de surface ($\log(g)$)	72
3.5.3	Métallicité ($[Fe/H]$)	74
3.5.4	Vitesse de rotation projetée ($v \sin(i)$)	76
3.6	Conclusion	78
4	Apport de la régression inverse par tranches (SIR)	93
4.1	Cas d'un problème linéaire	94
4.1.1	Régressions locales	95

4.1.2	k-plus proches voisins	96
4.1.3	Conclusion sur l'exemple linéaire	96
4.2	Cas d'un problème non-linéaire	96
4.2.1	Régressions locales	97
4.2.2	k-plus proches voisins	98
4.2.3	Sélection de directions pertinentes	98
4.3	Cas d'un problème complexe	101
4.4	Application sur des données synthétiques	101
4.5	Application aux données réelles	104
4.6	Conclusion sur l'utilisation de SIR	105
5	Étude comparative	117
5.1	Méthodes de référence	117
5.1.1	Moindres carrés partiels	117
5.1.2	MATISSE	119
5.2	Tables comparatives	120
	Conclusion	123
	Bibliographie	125
	Annexes	133

Introduction

Le développement de méthodes automatiques de traitement de données est, dans tous les domaines d'applications scientifiques, poussé par l'évolution technologique qui permet la création d'outils dédiés à la science de plus en plus performants. Cela passe par l'instrumentation, avec (Pennycook et al., 2006) par exemple, qui montre que la résolution des microscopes électroniques permet désormais de voir en dessous du dixième de nanomètre, les capacités de calcul (Henning, 2000) ou de stockage d'information (Orlov et al., 2004). Quand le manque de précision ou de données étaient, il y a quelques dizaines d'années, la cause principale de l'incompréhension d'un phénomène ou de l'imprécision d'une mesure, la capacité croissante d'acquisition de ces grandes quantités de données, tant due à la précision et la dynamique des mesures qu'à la multiplication des capteurs, soulève des problèmes différents.

En effet, disposant de beaucoup plus de données, et de données de plus en plus diverses, il est statistiquement plus probable que l'information recherchée, quelle qu'elle soit, s'exprime dans les données collectées. "Quelle qu'elle soit", effectivement, car quelle est-elle ? Disposer de ces grandes quantités de données nous ouvre de nouvelles portes et nous sommes tentés d'essayer d'accéder à l'information pertinente à l'aide d'outils statistiques.

Ces méthodes de traitement statistique de données peuvent notamment permettre d'extraire l'information sur une caractéristique d'un objet étudié à partir des données collectées, sans avoir précisément de modèle *a priori* du lien qui existe entre cette caractéristique et les données.

C'est dans ce contexte que les techniques de *machine learning* (Freitag, 2000) prennent de l'importance. Le *machine learning* regroupe des méthodes computationnelles à apprentissage, dont une partie repose sur l'emploi de méthodes statistiques pour permettre à un système automatisé de résoudre un problème de traitement de données. Ainsi, grâce à ce type de méthodes, on peut déduire, par apprentissage, à partir d'une base d'exemples connus, un lien entre des données acquises pour un objet d'étude et des caractéristiques associées à cet objet. Par exemple, en médecine, il peut s'agir de déterminer le lien statistique entre la caractéristique de risque d'incident cardiovasculaire d'un être humain en fonction de données telles que son poids, son âge, sa taille, son sexe etc.. L'objectif est de déterminer la valeur de la caractéristique associée à un nouveau jeu de données. Dans le cas de l'étude présentée dans ce mémoire, l'application étant issue d'une problématique de physique stellaire, les caractéristiques sont les paramètres stellaires fondamentaux (température effective, gravité de surface, métallicité et la vitesse de rotation projetée), et les données sont les spectres lumineux d'étoiles mesurés.

Le domaine de l'astrophysique a, comme les autres domaines scientifiques, aussi bénéficié de ces évolutions technologiques : amélioration des performances des CCD¹ (David Dussault, 2004), qui sont les transducteurs les plus utilisés en observation stellaire, amélioration de la qualité des optiques, réduction de l'encombrement du matériel spatial embarqué, taille des serveurs de données et capacité de calcul croissante des processeurs etc.. Cela a conduit, en physique stellaire, à agréger un grand nombre de spectres stellaires issus d'étoiles dont on souhaite déterminer le plus précisément possible les caractéristiques physiques. La question qui se pose est alors la suivante : peut-on (et si oui dans quelle mesure et avec quelle précision ?) avoir une méthode fiable et automatique permettant de déterminer les caractéristiques physiques et chimiques de ces étoiles à partir de la mesure de leurs spectres, et ce, en partant d'une base de données d'étoiles connues ou de spectres synthétiques ?

La problématique présentée s'intéresse donc à la détermination des paramètres stellaires fondamentaux, qui sont la température effective, la gravité de surface et la métallicité, à

1. Charge-Coupled Devices

partir de spectres à haute résolution. Au départ de cette thèse, la communauté astrophysique était déjà dotée d'outils pour effectuer ce type de travaux, comme : une régression basée sur l'analyse en composantes principales (Paletou et al., 2015a) ou l'algorithme développé pour estimer les paramètres stellaires fondamentaux à partir des données du RVS² de Gaia³ (Recio-Blanco et al., 2006). Du point de vue du domaine d'application, le but de cette thèse est d'apporter une alternative plus performante pour le traitement des données de spectroscopie stellaire.

Nous nous orienterons vers des méthodes de traitement statistique de données connues, utilisées par la communauté du traitement du signal, que nous adapterons au traitement des données de spectroscopie stellaire.

Le second objectif, implicite mais néanmoins important est l'optimisation de ces méthodes de traitement statistique compte-tenu de la particularité des données de type spectroscopiques. Nous savons que le lien entre les données (les spectres stellaires) et les paramètres est non-linéaire. Cependant la nature de ces non-linéarités est très complexe et souvent mal modélisée. Nous voulons donc exploiter le lien entre les données et les paramètres, sans connaître la nature des non-linéarités qui composent ce lien.

Nous appliquerons des méthodes de projections linéaires et proposerons plusieurs approches permettant la prise en compte des non-linéarités.

Dans le premier chapitre, nous nous attacherons à contextualiser la problématique du point de vue de l'application en retraçant l'évolution de la classification stellaire et des enjeux liés à l'estimation des paramètres stellaires fondamentaux. Le chapitre suivant présentera les méthodes de traitement statistique de données sur lesquelles nous avons axé nos efforts ainsi que des méthodes faisant référence, en termes de méthode de régression ou pour le traitement des données spectroscopiques, qui seront utilisées pour l'évaluation des résultats. Au chapitre 3, nous présenterons la régression basée sur l'analyse en composante principale, qui marque le point de départ du travail de thèse. Dans ce même chapitre seront ensuite présentées les méthodes de traitement alternatives que nous avons développées

2. *Radial Velocity Spectrometer*, Spectromètre équipant la mission Gaia

3. Mission lancée en 2013 pour répertorier et cartographier 1.7 milliard d'objets

pour la prise en compte des non-linéarités. Le chapitre 4 sera consacré à l'évaluation de l'apport de la régression inverse par tranches (SIR)⁴ en lieu et place de l'analyse en composantes principales (ACP)⁵ associé aux méthodes de traitement développées qui auront été présentés au chapitre 3. Enfin, nous présenterons dans le dernier chapitre une étude comparative des performances des méthodes proposées et des méthodes de référence sur des données d'étude, puis sur les données réelles.

4. *Sliced Inverse Regression*

5. Analyse en composantes principales

Chapitre 1

Contexte astrophysique

Nous cherchons à développer des méthodes de traitement de données propres au traitement d'un problème classique de physique stellaire : l'estimation des paramètres fondamentaux d'étoiles à partir de spectres lumineux. Il s'agit donc d'extraire l'information concernant la température effective, la gravité de surface et la "métallicité"¹ d'une étoile à partir d'un spectre mesuré. La température effective, notée T_{eff} , est une température caractéristique des couches externes de l'étoile. Elle est exprimée en Kelvin.

Pour des raisons de convention, nous utiliserons la grandeur $\log(g)$ le logarithme décimal de la gravité de surface, mesurée en dex².

Pour la métallicité, nous utiliserons l'abondance relative de fer notée $[Fe/H]$ comme indicateur. Elle sera mesurée en dex, et c'est une valeur relative à la métallicité du soleil³. Un quatrième paramètre retiendra notre attention : il s'agit de la vitesse de rotation projetée sur l'axe de visée. Bien que non fondamental, par rapport à l'étoile, ce paramètre a une grande influence sur les données que l'on mesure. En effet, cette vitesse de rotation projetée se traduit par effet Doppler en un élargissement des raies d'absorption. L'intégration de la

1. Quantité représentative de la composition chimique globale

2. Le dex est une unité logarithmique. Un intervalle de x dex représente une différence de 10^x de la grandeur de départ. Par exemple, la Terre a une accélération de la pesanteur de 9.81 m/s^2 donc 981 cm/s^2 (dans la mesure où le système d'unités "cgs" est encore très largement en vigueur en astrophysique) qui correspond à un $\log(g) \approx 3$.

3. La métallicité du soleil est nulle par convention.

lumière sur toute la surface visible de l'étoile somme les composantes de la lumière affectées d'un "blue-shift" (émise par la partie de l'étoile qui se rapproche) et les composantes affectées d'un "red-shift" (qui s'éloignent). Il convient donc de l'estimer avec la plus grande précision. Ce paramètre de vitesse projetée sera noté $v \sin(i)$ où i est l'angle d'inclinaison entre l'axe de visée et l'axe de rotation de l'étoile. Avant de présenter dans les chapitres ultérieurs, les méthodes de traitement statistique utilisées ainsi que nos développements, ce chapitre propose de présenter plus largement le problème de physique stellaire en retraçant son histoire puis en présentant quelques techniques actuellement appliquées pour résoudre ce problème.

1.1 Brève histoire de la classification spectrale

La spectroscopie stellaire débute au milieu du XIX^e siècle grâce aux premières observations du soleil effectuées par le physicien Joseph Van Fraunhofer en 1814. C'est cependant au britannique William Huggins que l'on attribue la paternité de cette science en raison du grand nombre d'observations qu'il a effectuées durant le siècle. Durant la seconde moitié du XIX^e siècle, c'est avec le travail de Giovanni Battista Donati (Donati, 1863) que l'on envisage la possibilité d'une classification spectrale à partir des observations spectroscopiques des étoiles.

Alors que la spectroscopie astrophysique se développe en Europe et aux États-Unis, avec G.B. Airy ou encore L.M. Rutherford⁴, il convient de citer aussi un astronome italien, Pietro Angelo Secchi, qui fut l'un des observateurs ayant produit le plus d'observations de son époque. Il peut revendiquer la classification de près de 4 000 étoiles aux alentours de 1870. Secchi dénombra trois classes en fonctions de la température : la première correspondant aux étoiles chaudes, la seconde aux étoiles de "type solaire" et la troisième correspondant aux étoiles froides. On peut observer l'impact de la modification de la température sur les spectres en figure 1.1. Par la suite, en 1877, Secchi proposera cinq classes distinctes.

4. On doit à ce dernier un autre moyen de classification spectrale basée sur les indices de couleur de nature photométrique et non pas spectroscopique.

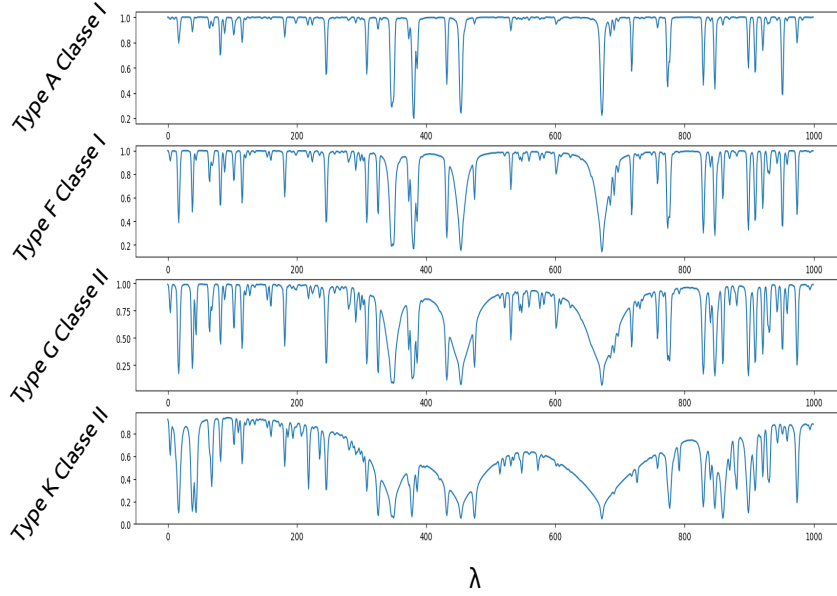


FIGURE 1.1 – Spectres synthétiques de types A, F, G et K (présentés au paragraphe suivant) appartenant aux deux premières classes de Secchi. Avec, en abscisse, les longueurs d’onde échantillonnées à pas constant, et en ordonnées, le flux lumineux pour chacune de ces longueurs d’onde. De haut en bas, ces spectres correspondent à des étoiles de plus en plus froides. Pour cet exemple les trois autres paramètres sont fixés à $\log(g) = 4.5$, $[Fe/H] = 0$, et $v \sin(i) = 0$.

À partir de 1885, au Harvard College Observatory (États-Unis) et sous l’impulsion de Charles Pickering notamment, est développée et mise en place une campagne de mesures spectroscopiques majeure. Celle-ci conduira, dans un premier temps, au “système de Draper”, en hommage à Henry Draper, pionnier de l’astrophotographie. En effet, à partir du

travail de classification de Pickering et de Fleming, des subdivisions aux classes de Secchi sont mises en évidence et un système de classification utilisant une série de seize lettres (incluant la classification des nébuleuses planétaires et “autres objets”) est alors adoptée.

On doit à Annie J. Cannon, dès 1901, le système de classes spectrales appelé “système de Harvard”. Ce système, affiné jusqu’aux alentours de 1912, est basé sur l’utilisation des lettres O, B, A, F, G, K et M, encore en vigueur aujourd’hui, qui servent à désigner les étoiles des plus chaudes aux plus froides. Chacune de ces classes est aussi subdivisée en sous-classes, au moyen d’un chiffre compris entre 0 et 9. Ainsi, par exemple, le Soleil est une étoile de type spectral G2.

Enfin, Le système de Morgan-Keenan-Kellman (Morgan et al., 1943), ajoute au système précédent une “classe de luminosité”. Cette dernière permet de distinguer les étoiles naines, des étoiles sous-géantes ou encore géantes, au moyen de l’ajout d’un chiffre en caractère romain au système de Harvard. Le Soleil est ainsi une étoile naine de type G2V alors que Bételgeuse sera une géante froide de type M2I, par exemple.

La figure 1.2 montre en abscisses les types spectraux et en ordonnées les classes de luminosités dans le diagramme dit de Hertzsprung-Russell. On peut ainsi y voir la répartition des étoiles au sein des différentes classes.

Au début du XX^e siècle, les théories physiques permettant de comprendre la nature et l’apparence des spectres n’existent pas encore. Il faudra attendre l’apparition de la mécanique quantique, puis de la physique atomique et moléculaire, aux alentours de 1920-30, pour être en mesure d’analyser ces observations de façon non-empirique, avec par exemple (Saha, 1921). Une autre difficulté à cette époque réside dans le fait que, c’est “à la main” ou plutôt “à l’œil”, et à partir d’une inspection détaillée et fastidieuse des spectres enregistrés sur des plaques photographiques, que s’effectuaient les identifications des objets dans telle ou telle classe. La classification reposait sur l’existence et/ou l’intensité de certaines raies spectrales, et la comparaison avec les spectres des objets déjà classifiés. La figure 1.3 montre la variation de l’intensité des raies spectrales caractéristiques en fonction du type spectral, issues des premiers modèles théoriques du XX^e siècle.

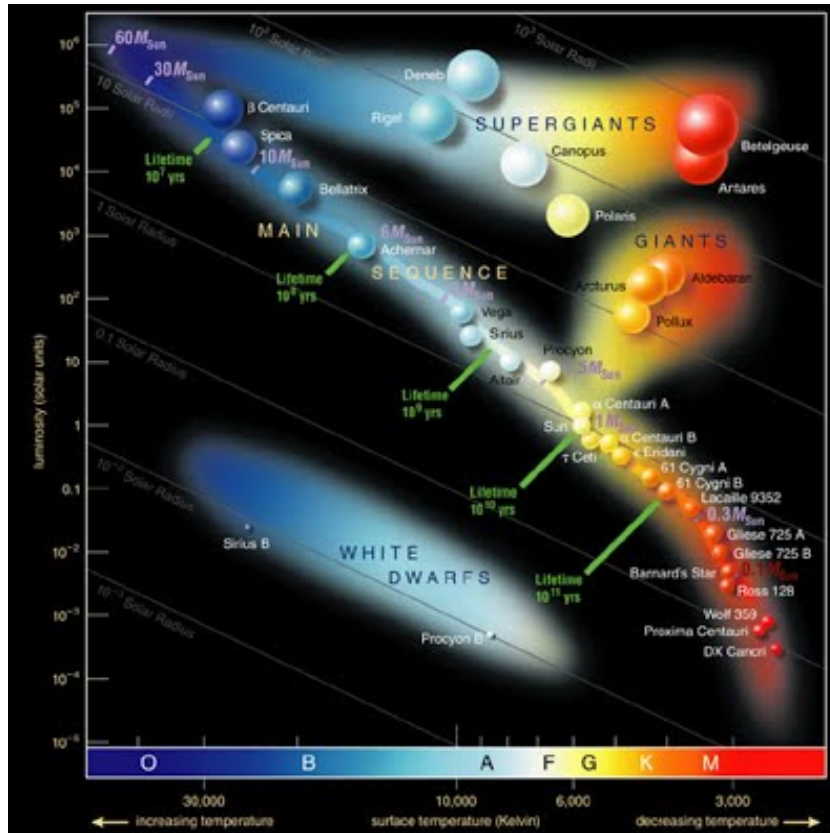


FIGURE 1.2 – Diagramme de Hertzsprung-Russell avec les classes de température du système de Harvard en abscisse, et les classes de luminosités apportées par le système Morgan-Keenan-Kellman. Les ordonnées portent les classes de luminosités. Les lignes diagonales représentent des lignes d'iso-gravité. On peut observer que les étoiles ne sont pas uniformément réparties sur cet espace. On peut les classer en trois catégories. La première catégorie est la séquence principale allant du haut gauche au bas droit du graphique. C'est dans cette catégorie que se trouve notre soleil. Les deux autres catégories sont les géantes et les super géantes d'une part, et les naines blanches d'autre part.

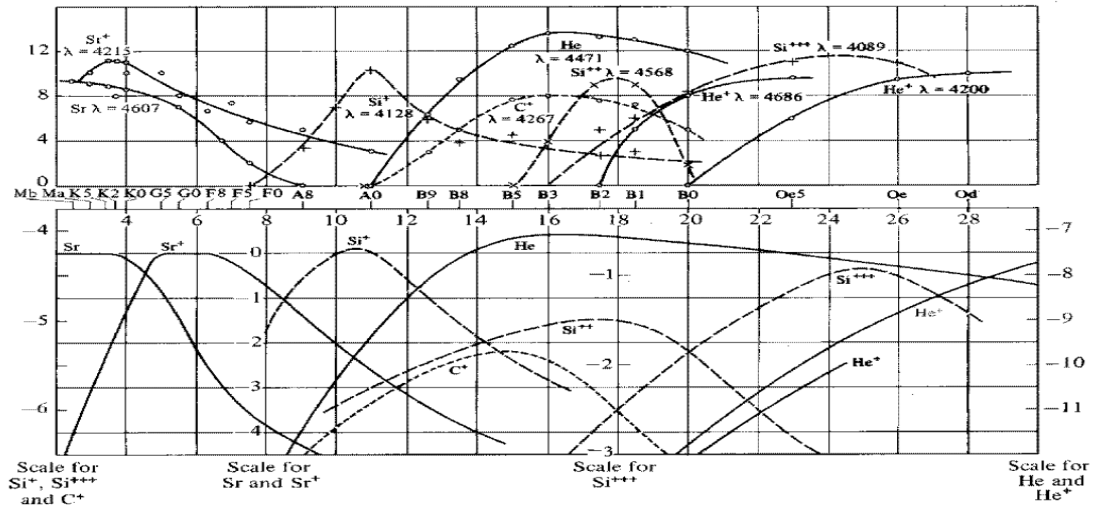


FIGURE 1.3 – Évolution de l'intensité des différentes raies spectrales en fonction de la température. Chaque courbe représente l'intensité d'une raie caractéristique associée à un élément chimique permettant de déterminer en abscisse le type spectral (d'après Payne, 1924).

1.2 Classification spectrale et paramètres stellaires fondamentaux

Dans les années 1970-80, la disponibilité et les capacités de calcul suffisantes des ordinateurs permettent le début des travaux de “classification automatique” des spectres stellaires. Il s'agira d'aller progressivement au-delà de la classification de Morgan-Keenan-Kellman (MKK) et de déterminer, plus finement donc, les “paramètres stellaires fondamentaux”. Cela est rendu possible grâce à l'essor de la modélisation numérique, et en particulier de la possibilité de modéliser des atmosphères stellaires, et les spectres associés, de façon assez réaliste, avec notamment, les codes pionniers de modélisation d'atmosphères stellaires comme Atlas : (Kurucz, 1970) (cf. figure 1.4) ou MARCS : (Gustafsson et al., 1975).

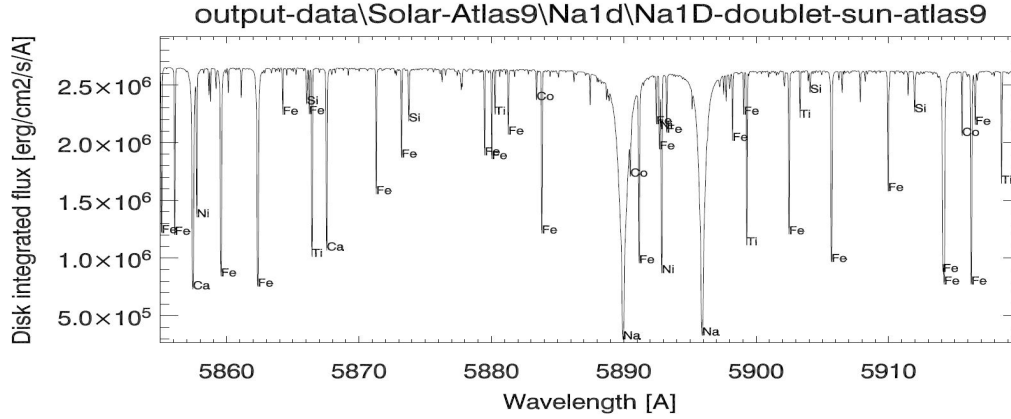


FIGURE 1.4 – Spectre synthétique issu d’un modèle d’atmosphère stellaire Atlas. Cette figure montre les raies d’absorption des éléments chimiques (cités dans le graphique au niveau de chaque raie) en fonction de la longueur d’onde.

Ces paramètres fondamentaux sont, pour rappel, la température effective, la gravité de surface et la métallicité. Ainsi, le Soleil, classé G2V dans le système Morgan-Keenan-Kelman, est aussi, du point de vue des paramètres fondamentaux, une étoile de température effective de l’ordre de 5780 K, de gravité de surface, $\log(g)$, d’environ 4.5 dex et de métallicité nulle (par construction), puisque la composition chimique du Soleil sert de référence aux autres étoiles.

Les premières tentatives de classification automatique reposent essentiellement sur des méthodes basées sur l’analyse en composantes principales (Deeming, 1964), (Whitney, 1983), les réseaux de neurones (von Hippel et al., 1994) voire la combinaison des deux techniques (Bailer-Jones et al., 1998), les méthodes de “minimum de distance”, (Katz et al., 1998) ou encore, dans une moindre mesure, une méthode de classification bayésienne, (Goebel et al., 1989).

Au milieu des années 1990, apparaissent les premières utilisations de réseaux de neurones, afin d’effectuer une classification automatique de distributions spectrales d’énergie stellaire (Gulati et al., 1994); (von Hippel et al., 1994). Dans un second temps, ces mêmes méthodes sont utilisées non plus pour classer les spectres (à basse résolution à cause des

ressources de calcul limitées) suivant le “système MKK”, mais plutôt en termes de valeurs des paramètres fondamentaux (Bailer-Jones, 2000).

À cette époque, la perspective de la mission spatiale Gaia, lancée en décembre 2013, commence déjà à prendre forme⁵ avec par exemple (Thévenin et al., 2003). En effet, il s’agira de traiter des informations, en particulier spectrales, au moyen de l’instrument *Radial Velocity Spectrometer* (RVS), pour plusieurs 10^5 objets plus ou moins intenses en termes de flux reçu par le capteur, et donc avec une grande disparité de rapport signal sur bruit (Bailer-Jones, 2013).

Nous sommes donc, depuis le début des années 2000, dans la perspective de cette grande campagne de mesures, mais aussi contemporains de la mise en œuvre et de l’analyse des données de nombreuses autres campagnes de mesures spectroscopiques⁶.

Ainsi, l’activité concernant les diverses approches de détermination des paramètres stellaires fondamentaux demeure très importante. Les développements de méthodes automatiques se basent sur des données obtenues à partir de calculs de spectres synthétiques, mais aussi à partir de nombreuses campagnes de mesures.

1.3 À propos des données astrophysiques étudiées

Dans un premier temps, nous évoquerons la détermination des paramètres fondamentaux, température effective, gravité de surface et métallicité, avec une méthode de minimum de distance suite à une analyse en composantes principales de spectres d’étoiles de type F-G-K (Paletou et al., 2015a). Pour ce faire nous avons utilisé deux jeux de spectres bien connus de la communauté astrophysique : ceux issus du spectrographe Elodie installé dans les années 1990 au télescope de 193 cm de diamètre (dit “193”) de l’Observatoire de Haute-Provence (Baranne, A. et al., 1996)⁷, et ceux issues du relevé “*Spectroscopic Survey of the Solar Neighbourhood*”, appelé S4N dans la suite du document, (Allende Prieto et al., 2004).

5. <http://sci.esa.int/gaia/>

6. Comme les SDSS (sdss.org), LAMOST (lamost.org) ou RAVE (rave-survey.org).

7. Cet instrument a permis la découverte de la première exoplanète autour d’une étoile, 51 Peg, en 1995

Les spectres Elodie sont des spectres à haute résolution ($R \approx 42000$)⁸ acquis sur une gamme de longueur d’onde visible allant de 390 à 680 nm. Ils sont aisément disponibles⁹ et ils ont déjà été largement analysés en termes de paramètres fondamentaux, grâce à un outil développé initialement par Katz et al. (1998).

Les spectres du S4N, que nous avons utilisés pour valider la méthode basée sur l’analyse en composantes principales (Paletou et al., 2015a), sont issus des spectrographes installés au télescope de 2.7 m d’ouverture du *Mc Donald Observatory* (États-Unis) et au télescope de 152 cm à La Silla (Chili). Dans chacun des cas, il s’agit de spectres de résolution $R \approx 50000$ pour une couverture spectrale de l’ordre de 360 à 1000 nm. Ils sont aussi largement disponibles ainsi que leur identification en termes de paramètres fondamentaux. C’est ainsi que nous avons pu nous baser sur des spectres réels et leurs identifications respectives, aussi bien pour Elodie que pour le S4N, afin de tester nos méthodes.

Nous avons aussi utilisé des spectres observés grâce aux spectropolarimètres de l’OMP Espadons et Narval (cf. figure 1.5), installés respectivement aux télescopes de 3.6 m CFHT (Canada-France-Hawaii-Telescope, États-Unis), et du Télescope Bernard Lyot de 2 m, au Pic du Midi de Bigorre (France). Dans les deux cas, puisqu’il s’agit d’instruments quasi-identiques, nous avons utilisé des spectres à résolution $R \approx 65000$, sur une bande spectrale couvrant de manière continue le domaine optique de 380 à 1000 nm, soit de l’ultra-violet proche, à l’infrarouge proche. Ces spectres sont distribués par le service d’observation de l’INSU sous la responsabilité de l’OMP, PolarBase¹⁰. La mise en œuvre d’outils de détermination des paramètres fondamentaux sur l’ensemble de cette ressource est à l’origine de notre projet de recherche de méthodes et donc à l’origine de cette thèse.

8. $R = \frac{\lambda}{\Delta\lambda}$ où λ est la longueur d’onde et $\Delta\lambda$ le pas d’échantillonnage en longueur d’onde

9. *Elodie spectral library* <http://atlas.obs-hp.fr/elodie/>; voir aussi : <https://arxiv.org/abs/astro-ph/0703658>

10. polarbase.irap.omp.eu

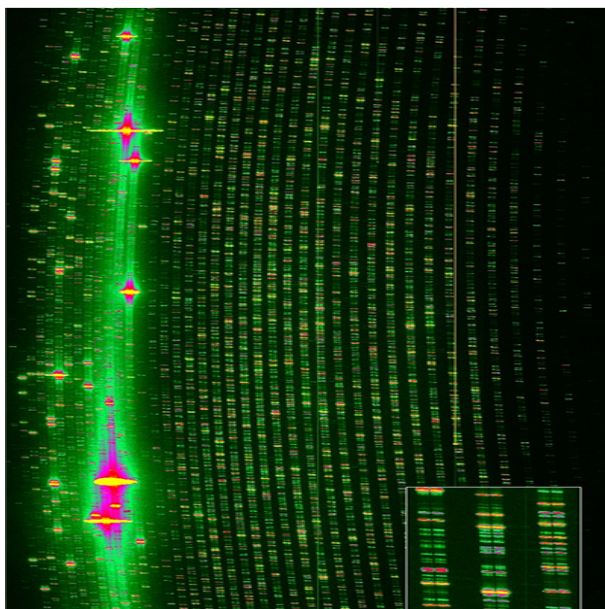


FIGURE 1.5 – Exemple de spectre brut typique issu d’Espadons ou Narval montrant la multiplicité des raies spectrales disponibles pour l’analyse.

Une deuxième étude dédiée aux étoiles froides de type naines M (Paletou et al., 2015b) a été réalisée sur des spectres observés, issus du spectrographe Harps de l’ESO (Chili). Ces spectres ont une résolution de l’ordre de $R \approx 115\,000$, et nous les avons utilisés dans une bande centrée autour du doublet du sodium neutre vers 589 nm. Pour cette étude, nous avons effectué des déterminations des paramètres fondamentaux au moyen de spectres synthétiques issus de calculs d’atmosphère stellaire de type “Atlas” (Kurucz, 2005) et de l’outil de synthèse spectrale “Synspec” de Hubeny & Lanz (1992).

Enfin, pour une étude ciblée sur la caractérisation de spectres d’étoiles chaudes de type A (Gebran et al., 2016) au moyen de spectres synthétiques Atlas également, nous avons aussi utilisé des données observationnelles issues du spectrographe, successeur de Elodie, Sophie, et installé lui aussi au télescope de 193 cm de l’OHP (S. Perruchot, 2008).

Nous avons vu au cours de ce chapitre que l'estimation de paramètres stellaires fondamentaux découle d'une importante histoire de classification spectrale en astrophysique. Nous souhaitons, par ces travaux de thèse, doter la communauté astrophysique d'un nouvel outil performant permettant l'estimation des trois paramètres stellaires fondamentaux que sont la température effective, la gravité de surface et la métallicité, mais aussi le paramètre de vitesse de rotation projetée. Partant des méthodes déjà développées nous étudierons des alternatives et proposerons une méthode performante pour le traitement des données spectroscopiques.

Chapitre 2

Méthodes de traitement statistique de données

Dans ce chapitre, nous nous intéresserons aux méthodes existantes pour le traitement des données. Ces données sont en particulier les spectres lumineux d'étoiles décrits au chapitre précédent. Les méthodes de traitement statistique de données doivent permettre une mise en forme des données pour une extraction plus aisée de l'information relative aux paramètres stellaires fondamentaux. Nous présenterons quatre méthodes permettant une réduction de la dimension de l'espace des données dans le but d'y appliquer une régression. Cette régression doit permettre d'obtenir le lien entre l'espace des données et celui du paramètre. L'objectif est de trouver le lien entre les spectres et les paramètres stellaires associés.

Les spectres sont nos vecteurs de données et les paramètres nos variables à expliquer. Nos données s'expriment dans un espace de très grande dimension et les variables à expliquer sont chacune portées par un espace de dimension 1. Les méthodes de régression nous permettront de trouver un lien statistique entre l'espace des données et les espaces des paramètres. En premier, vient une régression basée sur l'analyse en composantes principales ou ACP (Jolliffe, 1986). Il s'agit d'une méthode non-supervisée permettant d'effectuer une réduction de dimensionnalité (Une méthode non-supervisée est une méthode pour laquelle

aucune information extérieure aux données n'est ajoutée pour trouver la projection optimale). L'ACP peut être utilisée comme un pré-traitement à une régression, car l'ACP permet la détermination d'un espace de projection pour les données, qu'il faut ensuite lier à l'espace du paramètre. En suivant, nous présenterons la régression inverse par tranches, SIR (pour *Sliced Inverse Regression*) (Li, 1991). Cette méthode propose une réduction de dimensionnalité supervisée¹, et donc plus orientée vers un problème de régression. Ensuite, nous présenterons la méthode des moindres carrés partiels (PLS ; Tenenhaus (1998)) en tant que méthode de référence pour les problèmes de régression linéaire. Enfin, nous présenterons une méthode issue du domaine de l'astrophysique et développée dans le but de traiter le même type de données que celles que l'on souhaite traiter, MATISSE (Recio-Blanco et al., 2006). Ces deux dernières méthodes permettront de comparer, dans le dernier chapitre, les performances atteintes par les développements que nous proposerons concernant l'ACP et SIR.

2.1 Analyse en Composantes Principales

L'ACP est une méthode dont le but est de trouver, pour des données, la suite ordonnée des axes orthogonaux qui maximisent la variance des données projetées (Jolliffe, 1986). Autrement dit l'ACP cherche à hiérarchiser les directions de l'espace des données en fonction de l'inertie² qu'elles portent lors de la projection des données. On obtient donc par l'ACP, la direction qui maximise la variance des données, puis la direction orthogonale à la première qui maximise le même critère, et ainsi de suite jusqu'à avoir autant de directions que la dimension de l'espace. En ne gardant que les premières directions, il est possible de réaliser une réduction de dimensionnalité.

1. En opposition aux méthodes non-supervisées, les méthodes supervisées prennent en compte l'information propre à la variable recherchée, pour la construction du sous-espace de projection.

2. L'inertie d'un groupe de points est la moyenne des carrées des distances de ces points au centre de gravité du groupe (Saporta, 2011).

2.1.1 Principe

On considère la matrice des données $\mathbf{X}$ de dimensions $M \times N$. Ainsi N est la dimension de l'espace de départ est M le nombre d'individus. Par la suite, nous désignerons par "individu", un couple composé d'un vecteur de données et d'un jeu de variables à expliquer soit, du point de vue de l'application, un spectre et les paramètres fondamentaux associés à ce spectre.

L'ACP construit une matrice de projection $\mathbf{V}$ de telle sorte que le vecteur $\mathbf{X}^T \mathbf{V}$ soit de variance maximale. On parle ici de variance inter-individus et non de la variance des colonnes de la matrice qui représenterait une variance inter-axes (axes qui pour l'application sont les longueurs d'onde mesurées). Il s'agit donc de la variance des lignes de la matrice.

On peut présenter l'ACP sous la forme d'un problème d'optimisation visant systématiquement à rechercher la projection linéaire qui maximise la variance des données. Ainsi, nous pouvons écrire la variance des données projetées³ comme :

$$\mathbf{V} = \frac{1}{N} \sum_{i=1}^N [v_j^T (x_i - \bar{x})]^T [v_j^T (x_i - \bar{x})] \quad (2.1)$$

où l'indice i représente l'individu de la base de données $\mathbf{X}$ traitée, v_j est le vecteur de projection linéaire et V la variance après projection. On peut lier la variance après projection avec la variance-covariance totale de $\mathbf{X}$ dans l'espace de départ, Σ de la façon suivante :

$$\mathbf{V} = v_j^T \Sigma v_j \quad (2.2)$$

On cherche donc à maximiser $J(v_j)$, représentant la variance projetée sur le vecteur propre v_j , tel que :

$$J(v_j) = v_j^T \Sigma v_j \mid v_j^T v_j = 1 \quad (2.3)$$

La contrainte $v_j^T v_j = 1$ permet d'éviter, lors de la minimisation du critère, la solution triviale impliquant $v_j = 0$. On cherche donc une solution à :

$$\frac{\delta J(v_j)}{\delta v_j} = 0 \quad (2.4)$$

3. La projection correspond à l'application d'une transformation linéaire des données suivie de la troncature de la dimension de l'espace sur lequel elles sont représentées.

différente de $v_j = 0$. On peut ainsi écrire le problème :

$$\frac{\delta}{\delta v_j} v_j^T \Sigma v_j = 0 \mid v_j^T v_j = 1 \quad (2.5)$$

Grâce aux multiplicateurs de Lagrange, on peut alors écrire :

$$L(v_i, \lambda) = v_j^T \Sigma v_j + \lambda(1 - v_i^T v_i) \quad (2.6)$$

La solution du problème contraint s'écrit donc :

$$\frac{\delta}{\delta v_i} L(v_i, \lambda) = 2\Sigma v_j - 2\lambda v_i = 0 \quad (2.7)$$

et on obtient donc :

$$(\Sigma - \lambda \mathbf{I}_M)v = 0, \quad (2.8)$$

ce qui signifie que v est le vecteur propre associé à la valeur propre λ de Σ .

Autrement dit, la première colonne de $\mathbf{V}$, v_1 , est la combinaison linéaire des différentes composantes de l'espace initial qui maximise la variance des données $\mathbf{X}$ projetées sur v_1 . La seconde colonne de $\mathbf{V}$, v_2 , est le vecteur orthogonal à v_1 et maximise le résidu de la variance de $\mathbf{X}$ sous la contrainte $v_1^T v_2 = 0$. Chaque colonne est un vecteur orthogonal à tous les autres⁴, et on conserve le nombre de colonnes correspondant à la dimension de l'espace de projection souhaité.

L'ACP est optimale pour conserver la variance dans les données ; cependant la conservation de l'information relative à un paramètre dépendra de son influence sur la variance des données : plus un paramètre engendrera par sa seule modification de grandes différences entre deux vecteurs de données, plus l'ACP permettra d'en faire ressortir l'influence. Une solution pour déterminer la matrice $\mathbf{V}$ consiste en la décomposition en valeurs propres et vecteurs propres de la matrice de variance-covariance de $\mathbf{X}$, notée Σ . Cela revient à résoudre le problème de l'équation 2.8 (Jolliffe, 1986). La figure 2.1 montre sur un exemple simple les directions issues de l'ACP.

4. L'orthogonalité est une conséquence de la symétrie de la matrice Σ .

Composante initiale 2

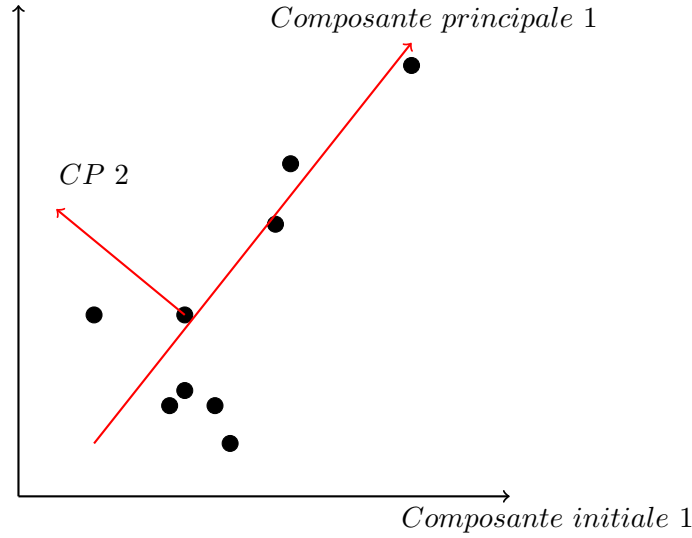


FIGURE 2.1 – Illustration de l’analyse en composantes principales. Ici, les données représentées par les points noirs sont dans un espace à deux dimensions matérialisé par les axes noirs. L’analyse en composantes principales permet d’obtenir la nouvelle représentation de l’espace matérialisé par les axes rouges (CP2 pour composante principale 2). La projection orthogonale des individus sur le plus grand des axes rouges permet une réduction de dimensionnalité, vers un espace de dimension 1, qui conserve un maximum de variance dans les données.

2.1.2 Mise en œuvre

La première étape dans la mise en œuvre de l’ACP est le calcul de la matrice Σ , matrice de variance-covariance des lignes de $\mathbf{X}$, où $\mathbf{X}$ la matrice des données :

$$\Sigma_{i,j} = \frac{1}{M} \sum_{k=1}^M (\mathbf{X}_{k,i} - \overline{\mathbf{X}}_i)(\mathbf{X}_{k,j} - \overline{\mathbf{X}}_j) , \quad (2.9)$$

où N est la dimension de l’espace initial, $\mathbf{X}_k$ est la k^e colonne contenant tous les individus x projetés sur la k^e composante de cet espace initial et $\overline{\mathbf{X}}_i$ et $\overline{\mathbf{X}}_j$ sont les valeurs moyennes

des colonnes i et j .

L'ACP permet (dans un cadre de réduction de dimensionnalité) d'obtenir un sous-espace décrit par les vecteurs propres associés aux plus grandes valeurs propres de Σ ⁵.

La matrice de projection $\mathbf{V}$ est de dimension $N \times P$ et chaque colonne est l'un des P vecteurs propres sélectionnés, où P est la dimension de l'espace sur lequel on souhaite projeter les données. La dimension de l'espace initial des données est notée N . Les individus x sont projetés vers le sous-espace donné par l'ACP de la manière suivante :

$$x_p = (x^T - \bar{\mathbf{X}})\mathbf{V} , \quad (2.10)$$

où x est un individu de la base de données représenté dans l'espace initial, et $\bar{\mathbf{X}}$ est le centre de gravité du nuage de points décrit par les individus dans l'espace initial. $\bar{\mathbf{X}}$ est calculé en faisant la moyenne empirique⁶ des lignes de la matrice $\mathbf{X}$. L'ACP nous permet la projection de tout vecteur de données vers le sous espace décrit par la matrice $\mathbf{V}$.

2.1.3 Propriétés

L'analyse en composantes principales présente des possibilités de débruitage (comme par exemple dans De La Hoz et al. (2015), où les auteurs filtrent des signaux grâce aux propriétés de l'ACP) dans le cas où les grandeurs acquises (mesurées) font partie d'un sous-espace de l'espace d'acquisition (de mesure), si les contributions du bruit⁷ sont statistiquement indépendantes. L'analyse de la décroissance des valeurs propres de la matrice de variance-covariance des données permet de déterminer le sous-espace comprenant toute l'information⁸ et le moins de bruit possible. L'ACP permet de présenter les données en maximisant la variance. En tant que méthode non-supervisée, elle ne prend pas en compte les variations du paramètre recherché.

L'analyse en composantes principales apporte aussi un intérêt en termes de compression du signal (Babenko et al., 2014). Du fait que l'on puisse garder l'information pertinente

5. Chaque valeur propre de Σ représente la variance des données portée par le vecteur propre correspondant.

6. La moyenne empirique correspond au moment statistique d'ordre 1 ou moyenne arithmétique

7. On entend par bruit la contribution des données indépendante de l'information que l'on recherche

8. L'information est ici la contribution du paramètre recherché dans les données

sur un sous-espace de dimension plus faible, l'information est portée par une quantité plus faible de données.

Nous pouvons utiliser l'ACP pour réduire la dimensionnalité des données et utiliser ses caractéristiques de débruitage. Cela nous permettra d'appliquer ensuite une méthode de régression depuis l'espace des données projetées vers l'espace du paramètre.

2.2 Régression inverse par tranches (SIR)

La régression inverse par tranche, présenté originellement par Li (1991), permet la recherche d'un sous-espace cohérent avec la variable que l'on souhaite expliquer (le paramètre). SIR permet d'utiliser des données recueillies sur un grand nombre d'individus pour retrouver un lien statistique vers une caractéristique de l'individu : la valeur d'un paramètre le caractérisant. En fait, on apporte, par rapport à l'analyse en composantes principales, un *a priori* sur ce que l'on recherche comme information dans les données. Là où l'ACP optimise la représentation des données pour "elles-mêmes" sur le principe du maximum de variance (sans considérer la valeur du paramètre que l'on recherche), SIR optimise la représentation des données pour une variable à expliquer bien précise. En effet, si le paramètre a un impact mineur sur les données, une méthode non-supervisée comme l'ACP ne permet pas de le mettre en valeur, alors que SIR en tant que méthode supervisée optimisera la projection pour ce paramètre en particulier. On peut déjà déduire que SIR, étant une méthode supervisée, impose un espace de projection différent pour chaque caractéristique. Lorsque l'on utilise l'ACP pour projeter des données en vue d'estimer la valeur d'un paramètre, il n'est pas nécessaire de reprendre le processus de création du sous-espace de projection lorsque le paramètre recherché est différent. SIR étant une méthode supervisée, il est nécessaire de créer un sous-espace propre à chaque paramètre.

2.2.1 Principe

Considérons un vecteur de données x , et une variable à expliquer (ou paramètre) y . SIR recherche les directions formant le sous-espace de projection le plus “cohérent”⁹ avec les variations du paramètre considéré. La question de comment la cohérence est évaluée est intéressante : il s’agit de maximiser une corrélation de manière indirecte¹⁰. La démarche est donc la suivante. On va tout d’abord s’intéresser aux variations de la valeur du paramètre pour les différents individus. C’est ainsi que l’on va subdiviser la base de données en H tranches en suivant la variation de la valeur du paramètre. D’après Li (1991), ces tranches sont voulues comme contenant chacune le même nombre d’individus. Ainsi, à l’intérieur d’une tranche les valeurs prises par le paramètre sont toutes comprises dans un intervalle de valeurs et aucun individu extérieur à cette tranche ne voit la valeur de son paramètre comprise dans cet intervalle. Une fois la base divisée, SIR va rechercher le sous-espace de projection linéaire, depuis l’espace des données, qui maximise la variance inter-tranches (chaque tranche étant représentée par son centre de gravité) tout en gardant une variance totale normalisée. Comme il n’y a pas de recouvrement entre les tranches, cela revient à chercher le sous-espace qui sépare au mieux les tranches en rapprochant au maximum les individus au sein d’une même tranche.

On peut y voir des similarités avec l’analyse factorielle discriminante (McLachlan, 2004) dans un contexte de classification, à ceci près que, dans notre cas, l’aspect continu des valeurs de paramètres fait que l’ordonnancement des valeurs prises par le paramètre a une signification. Pour le problème qui nous intéresse, nous avons des valeurs de paramètres que l’on peut ordonner de manière hiérarchique : 5000 K est plus froid que 5100 K. Alors que dans un contexte de classification de formes géométriques, par exemple, on ne peut pas dire qu’un carré est au-dessus ou en-dessous d’un cercle.

9. La cohérence ici est une corrélation entre l’évolution des données projetées et les variations du paramètre.

10. La corrélation est obtenue par la maximisation de la variance entre les individus ayant entre eux des valeurs éloignées de la caractéristique étudiée tout en normalisant la variance totale.

Nous appellerons $\mathbf{X}_h$ les individus de la base compris dans la tranche h . Pour chacune des tranches on calculera le centre de gravité $\bar{m}_h$:

$$\bar{m}_h = \frac{1}{n_h} \sum_{\mathbf{X}_i \in h} \mathbf{X}_i , \quad (2.11)$$

où n_h est le nombre d'individus dans la tranche h . Ainsi nous obtenons la matrice globale $\mathbf{X}_H$ par concaténation des $\bar{m}_h$. Chaque $\bar{m}_h$ est une ligne de la matrice $\mathbf{X}_H$. On peut ainsi calculer la matrice de variance-covariance inter-tranches $\mathbf{\Gamma}$:

$$\mathbf{\Gamma} = \frac{1}{H} (\mathbf{X}_H - \overline{\mathbf{X}_H})^T (\mathbf{X}_H - \overline{\mathbf{X}_H}) . \quad (2.12)$$

On souhaite ici maximiser la variance inter-tranches tout en normalisant la variance totale, là où l'ACP maximisait la variance globale. On écrit donc la variance inter-tranches $\mathbf{V}_{inter}$ des données projetées comme :

$$\mathbf{V}_{inter} = v_j^T \mathbf{\Gamma} v_j , \quad (2.13)$$

et la variance-covariance totale $\mathbf{V}_{tot}$ des données projetées :

$$\mathbf{V}_{tot} = v_j^T \mathbf{\Sigma} v_j . \quad (2.14)$$

On peut opérer la maximisation sous contrainte de la variance inter-tranches en maximisant le rapport $\mathbf{V}_{inter}/\mathbf{V}_{tot}$. Le critère à maximiser s'écrit :

$$J(v_j) = \frac{v_j^T \mathbf{\Gamma} v_j}{v_j^T \mathbf{\Sigma} v_j} | v_j^T v_j = 1 \quad (2.15)$$

(toujours sous contrainte de norme unité pour v_j). Or on sait que :

$$\frac{\delta}{\delta(x)} \left(\frac{a(x)}{b(x)} \right) = \frac{a'b - b'a}{b^2} . \quad (2.16)$$

On peut donc réécrire le critère :

$$\frac{\delta}{\delta v_j} J(v_j) = \frac{(v^T \mathbf{\Sigma} v) \frac{\delta}{\delta v_j} (v^T \mathbf{\Gamma} v) - (v^T \mathbf{\Gamma} v) \frac{\delta}{\delta v_j} (v^T \mathbf{\Sigma} v)}{(v^T \mathbf{\Sigma} v)^2} . \quad (2.17)$$

On cherche v_j tel que $\frac{\delta}{\delta v_j} J(v_j) = 0$, donc tel que :

$$(v^T \mathbf{\Sigma} v) \frac{\delta}{\delta v_j} (v^T \mathbf{\Gamma} v) - (v^T \mathbf{\Gamma} v) \frac{\delta}{\delta v_j} (v^T \mathbf{\Sigma} v) = 0 , \quad (2.18)$$

et comme Σ et Γ sont symétriques, on a :

$$(v^T \Sigma v) 2\Gamma v = (v^T \Gamma v) 2\Sigma v , \quad (2.19)$$

$$\Gamma v = \frac{(v^T \Gamma v)}{(v^T \Sigma v)} \Sigma v . \quad (2.20)$$

En posant $\lambda_i = (v^T \Gamma v)/(v^T \Sigma v)$. on obtient l'expression à résoudre :

$$(\Sigma^{-1} \Gamma - \lambda_i \mathbf{I}) v_i = 0 \quad (2.21)$$

Les vecteurs propres associés aux plus grandes valeurs propres de $\Sigma^{-1} \Gamma$ seront les directions construisant le sous-espace de SIR. La matrice Σ est la matrice de variance-covariance globale vue dans le paragraphe 2.1. Les vecteurs qui vont permettre de construire le sous-espace sont les solutions de l'expression :

$$(\Sigma^{-1} \Gamma - \lambda \mathbf{I}_M) v = 0 .$$

On peut considérer que l'ordre des tranches représente la variation du paramètre dans l'espace des données. SIR cherche la projection qui sépare au mieux ces tranches. La sensibilité de la méthode aux non-linéarités dans l'expression du paramètre dépend de la taille des tranches que l'on choisit. Si les tranches sont trop grandes, la méthode risque de ne pas être sensible à des non-linéarités représentatives de l'expression du paramètre dans l'espace des données. Si l'on considère des tranches contenant trop peu d'individus, les centres de gravité risquent de ne pas être représentatifs de l'expression du paramètre dans l'espace des données. SIR permet donc d'extraire les directions en fonction de l'expression du paramètre dans les données. Le sous-espace formé par ces directions est donc plus propice à l'application d'une régression que ne l'était celui de l'ACP.

2.2.2 Mise en œuvre

La mise en œuvre de la méthode peut être décrite de la manière suivante :

1. Calculer la matrice de variance-covariance globale Σ ;
2. Ordonner le vecteur contenant les valeurs du paramètre y . On l'appellera alors y_{tri} ;

3. Diviser le vecteur y_{tri} en H tranches, sans recouvrement ;
4. Pour chaque tranche h , calculer le centre de gravité $\bar{m}_h$ comme vu à l'équation 2.11 ;
5. Construire la matrice $\mathbf{\Gamma}$ (eq. 2.12) ;
6. Construire la matrice de passage $\mathbf{V}$ de sorte que les colonnes de $\mathbf{V}$ soient les vecteurs propres associés aux plus grandes valeurs propres de $\mathbf{\Sigma}^{-1}\mathbf{\Gamma}$.

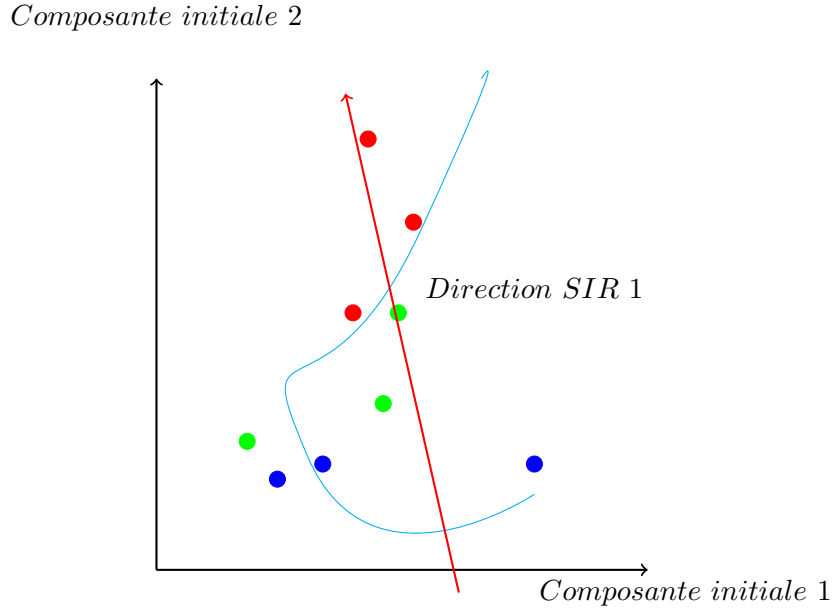


FIGURE 2.2 – Illustration de *Sliced Inverse Regression*. Ici les individus sont toujours représentés par les points dans un espace à deux dimensions matérialisé par les axes noirs. Mais contrairement à la figure 2.1, on considère la valeur d'un paramètre qui nous fait répartir en tranches les individus suivant la valeur prise par ce paramètre (correspondant à la projection des points sur la courbe bleue). Ces tranches sont visibles par les différentes couleurs de points. La méthode SIR permet de trouver l'axe rouge comme étant la direction la plus cohérente avec les variations de la variable à expliquer.

2.2.3 Propriétés

Comme expliqué au paragraphe 2.2.1, SIR nécessite un sous-espace différent (adapté) pour chaque paramètre (cf. figure 2.2.2). Ainsi, si l'on souhaite estimer les valeurs de plusieurs paramètres à partir des données, il sera nécessaire de construire plusieurs sous-espaces. La faiblesse de la méthode est sa sensibilité au conditionnement de Σ . Les directions de projections sont calculées en passant par l'inversion de la matrice Σ . Si celle-ci est mal conditionnée, les directions ne seront plus du tout représentatives des données.

SIR a quand même l'avantage certain, en tant que méthode supervisée, de tenir compte de la connaissance des valeurs du paramètre de la base de référence lors de la création de l'espace. Cet espace devrait donc être par construction plus pertinent. La mise en œuvre est néanmoins rendue plus complexe par le choix de la taille des tranches. Une façon de s'en affranchir est de procéder par un protocole de validation croisée¹¹ pour chaque type de données et avoir ainsi le nombre optimal de tranches (au sens de l'erreur d'estimation minimale) systématiquement.

2.2.4 Conditionnement de la matrice de covariance

Au paragraphe 2.2.3, nous avons mentionné le conditionnement de la matrice Σ et de son influence sur le fonctionnement de la méthode. En effet, dès qu'un problème implique l'inversion d'une matrice, le conditionnement de celle-ci peut avoir un énorme impact sur la sensibilité des résultats à une petite perturbation. Le conditionnement $K(\mathbf{A})$ d'une matrice $\mathbf{A}$ s'exprime comme suit :

$$K(\mathbf{A}) = \|\mathbf{A}\| \|\mathbf{A}^{-1}\| . \quad (2.22)$$

Ce conditionnement détermine la propagation de l'erreur pendant l'inversion d'un système (Petit & Maillet, 2008). Pour le cas d'un système $\mathbf{A}x = b$, connaissant $\mathbf{A}$ et b , on peut obtenir x grâce à l'inverse de A :

$$x = \mathbf{A}^{-1}b . \quad (2.23)$$

11. La validation croisée consiste à faire un grand nombre d'essais avec une base de données connue, segmentée aléatoirement et différemment à chaque essai, pour garder au final le nombre de tranches qui permet d'obtenir les meilleurs résultats d'estimation en moyenne.

Si l'on appelle $\|\delta_b\|$ l'erreur commise sur b et $\|\delta_x\|$ l'erreur engendrée sur x , alors on peut calculer le lien grâce à $K(\mathbf{A})$, conditionnement de $\mathbf{A}$:

$$\delta_x \geq \delta_b * K(\mathbf{A}) . \quad (2.24)$$

Dans le cas de SIR, la détermination des directions de projections données vient des vecteurs propres d'une matrice que nous appellerons $\mathbf{S}$. Cette matrice $\mathbf{S}$ est issue du produit de $\mathbf{\Gamma}$ avec l'inverse de la matrice $\mathbf{\Sigma}$. Si l'on reprend le raisonnement précédent, on peut conclure qu'une "erreur" sur $\mathbf{\Gamma}$ se répercute sur $\mathbf{S}$ en étant multipliée par le conditionnement de $\mathbf{\Sigma}$. La matrice $\mathbf{S}$ risque dans le cas d'un mauvais conditionnement de $\mathbf{\Sigma}$ de ne plus porter l'information de $\mathbf{\Gamma}$. Dans ce cas, $\mathbf{S}$ ne représenterait plus correctement les données au sens où l'utilisation de SIR le voudrait. Pour améliorer le conditionnement de $\mathbf{\Sigma}$, il existe différentes approches de régularisation. Nous utiliserons la "troncature de la décomposition en valeurs singulières" (Kaipio & Somersalo, 2005) qui revient à une réduction de dimensionnalité par l'application d'une ACP *a priori* sur les données. Ce choix est guidé par la simplicité de cette méthode (troncature des valeurs singulières) comparée à une autre méthode de régularisation, mais aussi par la méconnaissance de la structure des données et du bruit rendant les méthodes de régularisation plus élaborées inefficaces. En réduisant l'espace de départ et en ne gardant que les composantes qui correspondent aux valeurs propres les plus grandes de $\mathbf{\Sigma}$, le conditionnement de celle-ci décroît.

2.3 Moindres carrés partiels

2.3.1 Principe

La méthode des "moindres carrés partiels" (*Partial Least Squares*) PLS est une méthode de régression présentée dans un contexte de régression multivariée appliquée, entre autres, en chimie (Wold et al., 1983). Les auteurs utilisent notamment PLS pour déduire la composition chimique d'échantillons à partir de données spectroscopiques.

PLS effectue la projection des données vers un sous-espace, comme le permettent l'ACP et SIR. PLS permet de projeter les données vers le sous-espace qui maximisera la corrélation

linéaire entre les données et le paramètre. Ensuite, la méthode effectue la projection de ce paramètre vers ce même espace, (Tenenhaus, 1998). La méthode recherche un espace représenté par $\mathbf{T}$ résumant au mieux les données $\mathbf{X}$ et le paramètre y de sorte que l'on puisse écrire :

$$y = \mathbf{T}\mathbf{Q}^T + \mathbf{E}_y, \quad (2.25)$$

$$\mathbf{X} = \mathbf{T}\mathbf{P}^T + \mathbf{E}_x. \quad (2.26)$$

où $\mathbf{P}^T$ et $\mathbf{Q}^T$ sont les matrices de projections permettant de projeter respectivement les données et le paramètre vers l'espace représenté par la matrice $\mathbf{T}$. La méthode cherchera donc à minimiser les deux erreurs $\mathbf{E}_x$ et $\mathbf{E}_y$.

La construction d'une composante t_i , colonne de la matrice $\mathbf{T}$, est issue d'une combinaison linéaire des données x :

$$t_i = \sum_j w_{ij} x_j \quad (2.27)$$

tel que :

$$w_{ij} = \frac{\text{cov}(x_j, y)}{\sqrt{\sum_j \text{cov}^2(x_j, y)}} \quad (2.28)$$

La concaténation des t_i permet alors d'obtenir la matrice $\mathbf{T}$.

2.3.2 Mise en œuvre

L'algorithme de construction de la matrice $\mathbf{T}$ se résume ainsi :

1. Initialisation : copie de la matrice $\mathbf{X}$ dans une matrice χ ;
2. Répéter jusqu'à atteindre la dimension souhaitée pour le sous-espace :
 - Calculer le vecteur $w^{(n)}$ corrélation des données et du paramètre à l'itération n :

$$w^{(n)} = \frac{\chi^T y}{\|\chi^T y\|} \quad (2.29)$$

- Calculer $t^{(n)}$ la composante de $\mathbf{T}$ à l'itération n :

$$t^{(n)} = \chi w^{(n)} \quad (2.30)$$

— Retrancher la contribution de la composante $t^{(n)}$ dans les données :

$$\chi = \chi - t^{(n)} w^{(n)T} T \quad (2.31)$$

3. On obtient alors $\mathbf{T} = \mathbf{X}\mathbf{W}$ et $\mathbf{Q} = (\mathbf{T}^T \mathbf{T})^{-1} \mathbf{T}^T y$;

Une fois les matrices $\mathbf{Q}$ et $\mathbf{W}$ connues, on peut estimer la valeur de y_i associée à n'importe quel vecteur x_i :

$$\hat{y}_i = x_i^T \mathbf{W} \mathbf{Q} \quad (2.32)$$

Dans le cas de PLS, la régression et la projection des données se font dans le même temps via les matrices $\mathbf{W}$ et $\mathbf{Q}$.

2.3.3 Propriétés

La régression PLS permet de trouver des variables latentes, c'est-à-dire celles qui sont à la fois corrélées à $\mathbf{X}$ et à y . La méthode recherche le sous-espace qui maximise la corrélation entre $\mathbf{X}$ et y . Dans le cas de PLS, le critère optimisé est une corrélation linéaire, ce qui pose un problème dans le cas des données spectroscopiques. Ces données ont en effet, des liens non-linéaires avérés avec les valeurs de paramètres recherchés.

La nature de ces liens non-linéaires n'étant pas ou mal modélisée, il nous est impossible de modéliser les non-linéarités pour les compenser. L'application de PLS se fera donc sous l'hypothèse contraignante que le lien entre données et paramètres est linéaire ou du moins qu'il peut être approché par une relation linéaire.

La méthode PLS bien qu'imposant une contrainte de linéarité est intéressante car c'est une méthode supervisée. Dans le cas où l'hypothèse linéaire serait respectée, PLS, étant optimisé pour ce cas, doit donner de meilleurs résultats que les méthodes précédentes.

2.4 MATISSE

L'algorithme *MATrix Inversion for Stellar SyntEsis* (MATISSE) (Recio-Blanco et al., 2006), présente pour nous un intérêt tout particulier parce qu'il a été développé pour l'estimation de paramètres stellaires à partir de données spectroscopiques. En effet, MATISSE

s'attache à la détermination de paramètres stellaires fondamentaux à partir d'un spectre mesuré, à l'aide d'une base de données de spectres connus.

2.4.1 Principe

La méthode est basée sur la maximisation d'une corrélation linéaire entre le spectre et la valeur du paramètre recherché. Ainsi, l'algorithme cherche un vecteur $B_\theta(\lambda)$ fonction des longueurs d'onde λ , qui permet de déterminer le paramètre stellaire θ lorsque le spectre observé $O(\lambda)$ est projeté sur $B_\theta(\lambda)$. Le vecteur $B_\theta(\lambda)$ est obtenu comme étant la combinaison linéaire optimale permettant de relier les spectres à leur paramètre.

2.4.2 Mise en œuvre

La première étape consiste à centrer les valeurs des paramètres et celles des données, c'est-à-dire retrancher à chaque valeur de paramètre la valeur moyenne du paramètre dans la base de données et à chaque spectre le spectre moyen, afin de pouvoir constituer un lien linéaire. Le fait de centrer les valeurs compense les phénomènes de biais. Pour cela, on soustraira à chaque paramètre et à chaque longueur d'onde son espérance estimée par une moyenne empirique issue de la base de données considérée. On construira ainsi :

$$B_\theta(\lambda) = \sum_i \alpha_i S_i(\lambda) , \quad (2.33)$$

où les α_i sont les coefficients de pondération optimaux des spectres $S_i(\lambda)$ et $B_\theta(\lambda)$ étant la combinaison linéaire pondérée des spectres de la base de données de sorte que :

$$\hat{\theta}_i = \sum_\lambda B_\theta(\lambda) S_i(\lambda) . \quad (2.34)$$

Lorsqu'un spectre de la base est projeté sur $B_\theta(\lambda)$, on doit obtenir l'estimateur (au sens de MATISSE) de son paramètre. Si l'on appelle c_{ij} le coefficient de corrélation :

$$c_{ij} = \sum_\lambda S_i(\lambda) S_j(\lambda) , \quad (2.35)$$

on peut réécrire ainsi l'équation 2.34 :

$$\hat{\theta}_i = \sum_j \alpha_j c_{ij} . \quad (2.36)$$

Les α_i maximisent globalement le critère R qui est la corrélation linéaire entre les θ_i (vraies valeurs de paramètres) et leurs estimateurs $\hat{\theta}_i$:

$$R = \frac{(\sum_i \hat{\theta}_i \theta_i)^2}{\sum_i \hat{\theta}_i^2} \quad (2.37)$$

D'après Recio-Blanco et al. (2006), R est maximisé pour :

$$\sum_k \left(\sum_i c_{ij} c_{ik} \right) \alpha_k = \sum_i c_{ij} \theta_i . \quad (2.38)$$

Les coefficients α_i sont donc les solutions de l'équation :

$$\sum_i c_{ki} \alpha_i = \theta_i . \quad (2.39)$$

On peut réécrire ceci sous forme matricielle où $\mathbf{C}$ est la matrice de variance-covariance :

$$\mathbf{C}\alpha = \theta . \quad (2.40)$$

Pour suivre la méthode d'application d'origine (Recio-Blanco et al., 2006), le mauvais conditionnement de la matrice $\mathbf{C}$ fait que le résultat est approché par un algorithme de Van Cittert (Bandžuch et al., 1997).

2.4.3 Propriétés

MATISSE propose un moyen de déterminer la combinaison linéaire optimale entre les spectres et la valeur du paramètre considéré. Pour cela, on peut faire un parallèle avec la méthode PLS présentée dans le paragraphe précédent. Même si MATISSE propose une autre approche, l'hypothèse de linéarité reste une contrainte forte. Comme dans le cas de PLS, cette méthode est optimale quand la contrainte de linéarité est respectée.

2.5 Conclusion

Dans ce chapitre nous avons décrit quatre méthodes de traitement statistique de données. Chacune permet de réaliser une régression des données depuis leur espace d'origine

vers les espaces des paramètres. Cependant, chacune optimise un critère différent. Là où certaines, PLS et MATISSE notamment, appliquent une régression complète depuis l'espace original des données vers l'espace des paramètres, les autres, telles que l'ACP et SIR sont utilisées comme des méthodes de projection. Il nous appartiendra dans les deux chapitres suivants de trouver la meilleure régression pour ces méthodes en proposant des outils permettant d'estimer les valeurs des paramètres fondamentaux à partir des données projetées.

Chapitre 3

Régression basée sur l'ACP

3.1 Description de l'utilisation de l'ACP pour des problèmes de régression

L'analyse en composantes principales (Jolliffe, 1986) s'inscrit comme la première étape, dans le problème de régression que nous étudions. On souhaite, grâce à celle-ci, représenter les données de la façon la plus pertinente pour permettre l'estimation des valeurs des paramètres. Ainsi, grâce à l'ACP, nous chercherons le sous-espace le plus pertinent possible pour une régression linéaire au sens des moindres carrés, ou d'autres approches permettant de tenir compte des non-linéarités comme une estimation au sens des "k-plus proches voisins" (k-PPV). Cette estimation par k-PPV consiste en une moyenne de la valeur du paramètre des individus, de la base de donnée de référence, dans un voisinage donné. Plus de détails sont présentés au paragraphe 3.1.2.

L'association de l'ACP et de l'estimation au sens des k-PPV a pu montrer de bons résultats pour l'estimation de paramètres stellaires fondamentaux (Paletou et al., 2015a).

3.1.1 Choix du nombre de composantes

Au chapitre 2, nous avons vu que l'ACP peut être utilisée dans un contexte de régression comme déterminant un espace de projection pour les données qui maximise leur variance.

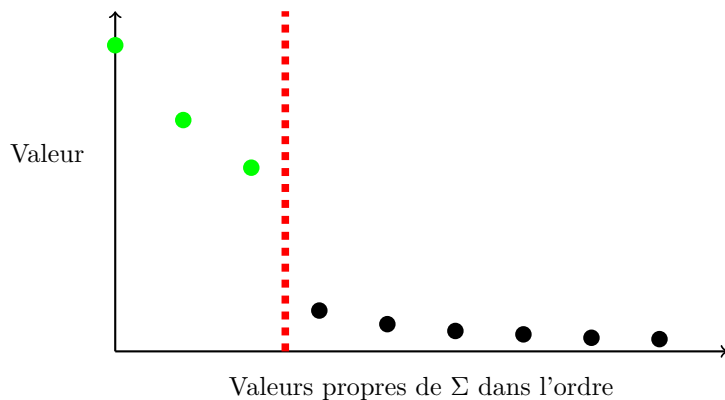


FIGURE 3.1 – Décroissance de valeurs propres typiques pour un faible niveau de bruit (dont la puissance est plusieurs ordres de grandeur inférieure à la puissance qui caractérise l’information pertinente : les différences intrinsèques entre les individus). La ligne en pointillé matérialise la “cassure”. Ici trois valeurs propres sont bien plus élevées que toutes les autres et l’espace de projection le plus pertinent sera donc celui décrit par les trois vecteurs propres associés.

La détermination de la dimension du sous-espace peut se faire de différentes manières. La décroissance des valeurs propres de la matrice de covariance des données Σ montre, suivant le niveau de bruit, une “cassure” (figure 3.1). Les premières valeurs propres sont beaucoup plus grandes que les suivantes. Les dernières valeurs propres, les plus “à droite”, sont presque au même niveau et bien plus faibles que celles “à gauche” de la “cassure”. Ce cas se produit lorsque le niveau de bruit est suffisamment faible, c’est-à-dire que la variance dans les données est issue des différences entre les individus et que le bruit qui vient s’ajouter aux données n’a pas beaucoup d’impact sur ces différences. On peut alors déterminer la dimension optimale de l’espace des données comme étant égale au nombre de valeurs propres qui sont grandes par rapport aux autres (cf figure 3.1), car ce sont celles associées aux vecteurs propres qui portent l’information indépendante du bruit.

La figure 3.1 montre que l’énergie (représentant la variance dans les données) est très concentrée dans les premières valeurs propres. Cela signifie qu’un espace de dimension

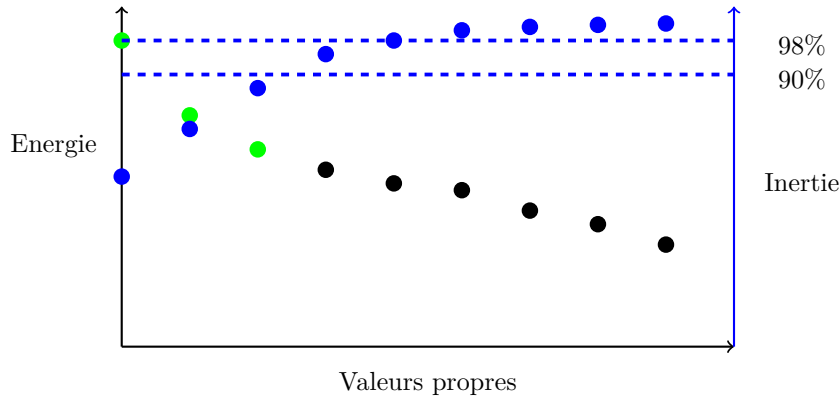


FIGURE 3.2 – En vert et noir, est représentée la décroissance de valeurs propres typiques pour un fort niveau de bruit. En bleu, on a le pourcentage d’inertie préservée lors de la projection des données vers le sous-espace associé à un nombre croissant de valeurs propres.

3 est suffisant pour représenter les données en conservant ainsi la variance intrinsèque des individus en éliminant une partie du bruit avec les composantes qui ne portent pas d’information. Le bruit, qui est dans ce cas “peu énergétique”, est porté en majeure partie par les vecteurs propres associés aux six dernières valeurs propres. Si l’on se place dans le cas d’un fort niveau de bruit, on peut se retrouver dans la situation de la figure 3.2.

Dans le cas d’un fort niveau de bruit que présente la figure 3.2, il est beaucoup plus difficile de savoir à partir de quelle composante l’expression de la puissance du bruit prend le dessus sur l’expression de la variance propre aux individus. Dans le cas où l’on a un fort niveau de bruit, on peut mesurer l’inertie¹ préservée et ainsi mettre un seuil en fonction du taux d’inertie que l’on accepte de perdre. Le problème est de trouver le bon taux d’inertie. Bien que ce ne soit pas optimal, fixer arbitrairement un seuil est possible. Si l’on connaît la puissance du bruit, on peut déterminer à partir de quelle composante la variance des données ne reflète plus que l’expression du bruit, et fixer un seuil de manière pertinente en se basant sur cet *a priori*.

Nous verrons que ce n’est pas le cas des données que l’on souhaite traiter. Nous opterons

1. Définie au chapitre 2, l’inertie revient au taux de variance conservée

plutôt pour un protocole de validation croisée². Le problème de cette approche est qu'elle dépend fortement de la méthode d'estimation qui suit. Un nombre de composantes optimal pour une estimation locale au sens des k -plus proches voisins n'a aucune raison d'être optimal dans le cadre d'une régression linéaire. Nous réappliquerons donc ce protocole pour chaque méthode d'estimation employée.

3.1.2 Choix de la méthode d'estimation

Une fois tous les individus projetés vers un sous-espace pertinent, il sera question de trouver le moyen de lier, de trouver la relation, entre l'espace des données projetées et l'espace du paramètre.

L'une des possibilités pour identifier un nouvel individu, c'est-à-dire d'estimer la valeur de son paramètre, est d'appliquer une régression linéaire permettant de passer dans l'espace du paramètre. Une autre méthode, qui a été utilisée en astrophysique avec Paletou et al. (2015a), consiste en la considération d'un voisinage, et à l'estimation de la valeur du paramètre par une moyenne calculée à partir des voisins considérés.

Régression linéaire

Une façon simple de lier deux espaces consiste à appliquer une régression au sens des moindres carrés de sorte qu'une fois projetées sur le régresseur³ les données expriment directement l'estimateur du paramètre. On considère les données $\mathbf{X}$ dont on souhaite connaître le lien avec le paramètre y suivant un modèle linéaire tel que $y = Xh + \epsilon$. Ici ϵ est une variable de bruit indépendante. Alors l'estimateur des moindres carrés de h s'écrit

$$\hat{h}_{MC} = (X^T X)^{-1} X^T y .$$

2. La validation croisée consiste à faire un grand nombre d'essais avec une base de données connue. À chaque essai, on extrait aléatoirement de la base de données un certain nombre d'individus qui serviront de test. Et à chaque essai, nous déterminons, le nombre optimal de composantes. La moyenne de ce nombre de composantes sur l'ensemble des essais sera la valeur retenue.

3. Vecteur de projection des données vers l'espace d'un paramètre

Il permet par la suite d'estimer la valeur de y pour un vecteur de données x en posant $\hat{y} = x^T \hat{h}$. La régression au sens des moindres carrés permet de trouver le régresseur qui minimisera la moyenne des distances quadratiques entre les données dans l'espace des données et leurs projections orthogonales sur le régresseur.

Estimation par k-plus proches voisins

L'approche par k-plus proches voisins (k-PPV) se base sur l'*a priori* que localement, dans l'espace des données projetées, les valeurs prises par le paramètre varient peu. On peut donc estimer la valeur du paramètre associé à un individu en faisant une moyenne (pondérée ou non) des valeurs prises par le paramètre chez ses voisins.

$$\hat{y}_{kppv} = \text{moy}(y_{\text{ref}} | x_{\text{ref}} \in \text{voisinage}) , \quad (3.1)$$

où "moy" est une fonction de moyenne que l'on peut appliquer suivant la métrique et les pondérations que l'on jugera pertinentes, et $x_{\text{ref}} \in \text{voisinage}$ correspond au sous-ensemble des individus de la base de référence inclus dans le voisinage. Ceux-ci peuvent être les k -plus proches comme dans la version originale avec un nombre de voisin invariant (Kataria & Singh, 2013), ou bien comme dans Paletou et al. (2015a) tous les voisins à une distance inférieure à une limite. L'avantage de cette dernière méthode est qu'elle est plus robuste aux non-linéarités. Il faut néanmoins définir le voisinage, soit en termes de nombre de voisins soit en termes de zone. Si le voisinage est trop grand, l'estimation sera très imprécise et si le voisinage est trop petit, l'estimation sera très sensible au bruit. Lorsque les données sont réparties de façon homogène, on peut alors définir une zone de voisinage fixe ou un nombre de voisins, mais si cette répartition n'est pas homogène il n'est pas garanti que tous les individus aient un voisin dans leur zone. Une manière de décider de la taille optimale pour le voisinage passe par la connaissance de la discernabilité⁴ des individus. Mais cela requiert d'avoir un *a priori* sur la qualité des mesures et sur la fiabilité des valeurs de paramètres de

4. Distance en deçà de laquelle deux individus ne peuvent être différenciés : par exemple si la base de référence échantillonne un paramètre avec un certain pas, il est impossible de différencier deux individus distants de moins de la moitié de ce pas, dans l'espace de ce paramètre.

la base de référence. À défaut, lorsqu'on ne dispose pas d'information suffisante à ce sujet, la solution de la validation croisée est envisageable. La validation croisée est efficace, mais ce procédé devient lourd si l'on doit le mettre en œuvre pour chaque étape (projection et estimation).

3.2 Particularité du traitement de données non-linéaires

Dans le cas d'un lien linéaire entre les données et le paramètre, une régression au sens des moindres carrés fonctionne très bien, mais ce n'est plus le cas lorsque le lien entre les données et le paramètre recherché devient non-linéaire. Par ailleurs, l'analyse en composantes principales étant une méthode de projection linéaire, elle ne permet donc pas à elle seule de prendre en compte les non-linéarités.

3.2.1 Régression non-linéaire

Les méthodes de régression non-linéaires classiquement utilisées dans la littérature ne sont pas utilisables pour le traitement des non-linéarités dont on ne connaît pas la nature. Les méthodes communes font appel à des *a priori* indisponibles dans le cas des données de spectroscopie stellaire. En effet, il est souvent nécessaire de connaître le type de non-linéarités auxquelles on fait face. Ainsi, on peut appliquer une régression aux données en les projetant sur une variété caractérisée par une fonction non-linéaire, mais dans ce cas il faut savoir quel type de fonction définit ladite variété. L'approche des courbes principales (*principal curves*), (Hastie & Stuetzle, 1989), permet de trouver une variété optimale pour appliquer une régression aux données, même si cette approche est limitée à des variétés 1D et repose sur la non-convergence de l'algorithme.

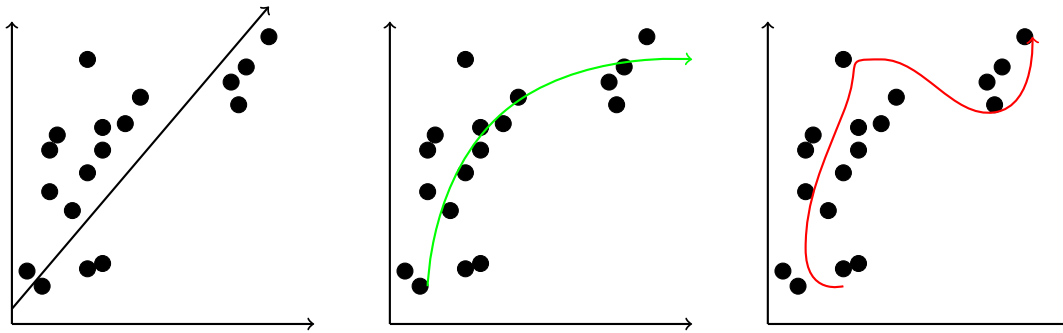


FIGURE 3.3 – Représentation de l’ACP à gauche, la *principal curve* idéale au centre et à droite dans le cas où l’on laisse l’algorithme aller trop loin et que la courbe suit le bruit. Les abscisses et ordonnées représentent les composantes de l’espace de départ.

Les courbes principales, illustrées à la figure 3.3, part de la première composante principale et cherche à partir de celle-ci, itérativement, une courbe qui minimise une erreur quadratique entre les points et leurs projections. Le problème est que l’algorithme converge vers une courbe qui passe par tous les points et ne représente plus les données mais le bruit. Pour l’estimation des valeurs des paramètres, il n’est pas nécessaire de connaître la variété en tout points de l’espace et on peut se contenter d’une régression linéaire qui approxime localement cette variété. C’est ce que nous présenterons par la suite.

3.2.2 Régression linéaire locale

Un moyen de pallier le problème des non-linéarités est de définir une zone plus restreinte de l’espace des données où l’hypothèse de linéarité est localement respectée autour de l’échantillon que l’on souhaite traiter. Un problème commun avec le critère d’arrêt des courbes principales consiste à définir la zone locale. Si la zone est trop petite, le régresseur est trop sensible au bruit ; si la localité est définie comme trop étendue, alors on perd complètement la sensibilité aux non-linéarités.

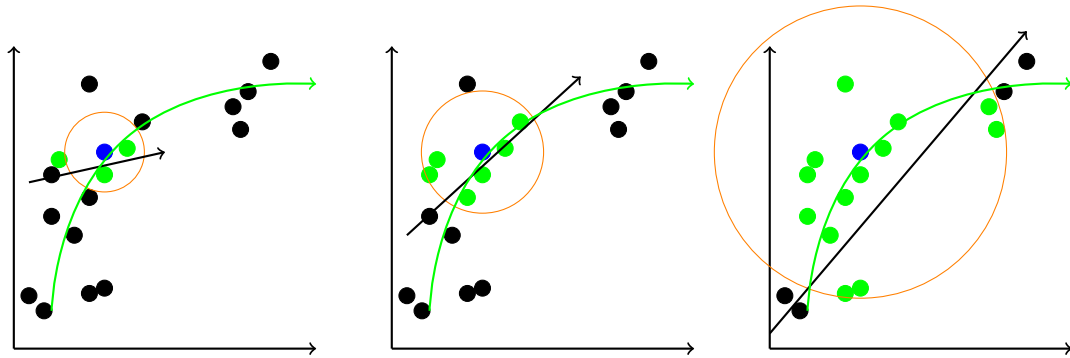


FIGURE 3.4 – Illustration des trois cas lors de la mise en place d’un régresseur local. Le but est d’approcher localement la courbe verte (liant les abscisses et les ordonnées) grâce aux directions de projection (flèches noires). Dans le cas de gauche, la zone considérée (celle à l’intérieur du cercle) est trop petite, dans le cas de droite, la zone considérée est trop grande, au centre, la zone permet d’avoir une droite qui, localement, autour du point bleu considéré, correspond à la courbe verte.

La figure 3.4 montre trois cas où l’on essaie d’approcher localement une variété par une droite. L’objectif est de se placer dans le cas du centre, là où la zone vérifie l’hypothèse de linéarité. Pour trouver le voisinage le plus efficace, on peut se demander ce qui se passe lorsque le voisinage est trop petit. Tant que l’on a une zone suffisamment grande, la corrélation que l’on recherche est plus grande que le bruit et la suppression de peu de points a peu d’impact. Dans le cas où le voisinage est trop petit, une modification légère dans la position d’un ou quelques points modifie beaucoup les caractéristiques du régresseur. On cherche donc le voisinage le plus petit possible, robuste à une faible modification. Une manière de trouver la taille optimale pour la zone est de définir le voisinage minimisant l’erreur de reconstruction, c’est-à-dire l’erreur moyenne entre les valeurs du paramètre des individus de la base de référence du voisinage, et leur estimation mutuelle par régression sur la direction obtenue.

3.2.3 Estimation basée sur les k-plus proches voisins

L'approche par la méthode des k-PPV ne fait pas appel à une méthode de régression linéaire. Dans le cadre d'une régression, l'hypothèse sous-jacente à l'emploi des k-PPV est que localement les valeurs du paramètre sont proches. Ainsi, cette méthode, par rapport aux régressions locales, traite tous les individus en définissant un voisinage suivant le même critère. On estime la valeur du paramètre de l'individu recherché⁵ comme étant la moyenne des valeurs du paramètre des voisins. Lorsque l'on applique ce type de méthodes, la définition du voisinage considéré est cruciale encore une fois. On souhaite donc définir un voisinage dont les valeurs du paramètre varient peu, voisinage dont le centre de gravité est suffisamment près de l'individu que l'on souhaite identifier pour que l'on puisse les considérer comme indiscernables. Il faut aussi faire attention de ne pas considérer trop peu de voisins, faute de quoi la méthode sera très sensible aux données aberrantes (ayant une valeur pour laquelle le bruit sort des statistiques qui le définissent).

3.2.4 Sélection de directions pertinentes

Ce paragraphe essaie d'apporter une solution en partant du principe que le sous-espace obtenu par la méthode de projection n'est pas partout le plus pertinent (quelle que soit la zone de l'espace où la projection est appliquée). Dans la mesure où le lien entre le paramètre recherché et les données serait non-linéaire, on espère par le biais de la méthode de projection (ici l'ACP) trouver le sous-espace de dimension la plus petite contenant la variété 1D correspondant à la variation des valeurs du paramètre. Mais il n'est pas évident que cette variété s'exprime sur tous les axes en tout point de l'espace. Il est bien possible que dans certains cas, pour certaines valeurs du paramètre, cette variété soit colinéaire (ou quasi colinéaire) à l'une des directions de l'espace, ou située dans un hyper-plan du sous-espace. Dans ce cas, considérer la totalité de l'espace pour l'estimation serait une erreur car on ferait entrer plus de bruit dans l'estimation. Un autre cas serait celui où pour une réalisation de x , un vecteur de données, la projection sur le sous-espace donne plusieurs

5. L'individu pour lequel on ne connaît que la valeur du vecteur de données mais pas les valeurs prises par les paramètres.

valeurs probables pour y . Ainsi la projection sur certaines directions va tendre à montrer que l'estimateur de y peut prendre deux valeurs $Y1$ ou $Y2$ différentes, avec chacune des probabilités comparables, là où d'autres directions ne seront pas ambiguës⁶. Dans ces zones du sous-espace, ignorer les solutions ambiguës permet de réduire l'erreur d'estimation. L'exemple simple de la figure 3.5 montre une réalisation de x dont la projection sur la première direction $x1$ a abouti à la valeur $X1$ ⁷ et la projection sur la seconde direction $x2$ a donnée la valeur $X2$. Or, lorsque l'on trace la densité de probabilité conditionnelle à cette réalisation, on observe deux modes espacés sur la direction 1 et un seul mode sur la direction 2.

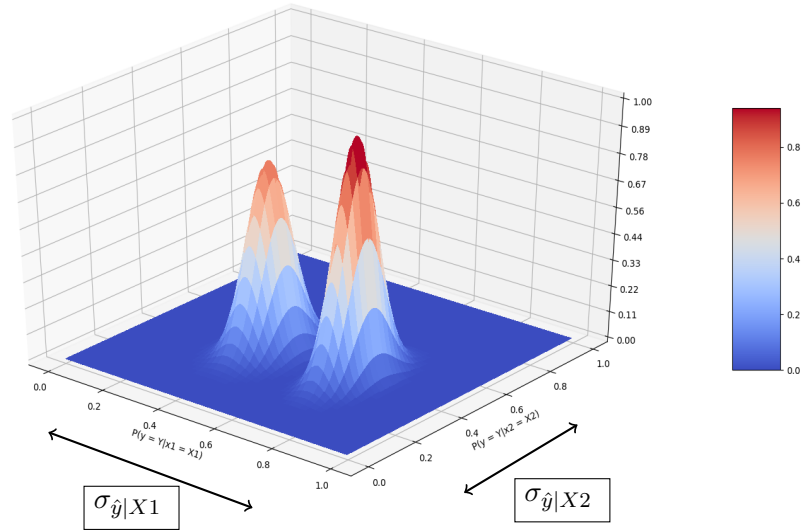


FIGURE 3.5 – Densité normalisée de $y = Y$ sachant une réalisation particulière de X . On observe que la valeur de y présente deux solutions suivant $x1$ et une seule solution suivant $x2$. Dans ce cas ne pas considérer $x1$ lève l'ambiguïté.

6. Les solutions ambiguës sont celles qui sont multiples (quasi-équiprobable) et très différentes là où nous cherchons une solution unique

7. $X1$ est une réalisation de la variable $x1$

Employer, dans cet exemple, l'intégralité du sous-espace ne permettrait pas d'obtenir une estimation pertinente de la valeur de y . En effet, une estimation par moyenne empirique aboutirait à un résultat autour de 0.5 (moyenne entre 0.4 et 0.6) qui est un résultat très peu probable quelle que soit la direction que l'on regarde. Le maximum de vraisemblance donnerait aussi un résultat ambigu, 0.8 ou 0.4. Ce que nous proposons, c'est de conserver la combinaison de directions (le sous-espace, et dans cet exemple, la direction) qui minimise la variance sur y dans l'espace de projection. Ainsi, en regardant la variance sur y des plus proches voisins suivant la projection choisie, on en vient à déterminer que :

$$\sigma_{\hat{y}|X2} < \sigma_{\hat{y}|X1} , \quad (3.2)$$

où $\sigma_{\hat{y}|X2}$ et $\sigma_{\hat{y}|X1}$ sont les variances des valeurs du paramètre sachant respectivement les réalisations $X1$ et $X2$ résultats des projections des données sur les directions $x1$ et $x2$.

Ainsi, nous pouvons appliquer indifféremment une méthode des k-PPV ou de régression linéaire, mais sur un sous-espace différent pour chaque individu. Le sous-espace optimal choisi pour chaque individu sera celui qui minimise la variance empirique sur y dans le voisinage de celui-ci.

On peut opposer à cette approche qu'il aurait sans doute été plus optimal de créer un sous-espace optimisé pour chaque individu en appliquant la méthode de projection (par exemple une ACP) sur le voisinage. Mais cela réduit le nombre d'individus utilisés pour le calcul du sous-espace et pose des problèmes de stabilité aux méthodes. En effet, les matrices de données se retrouvent rapidement avec beaucoup moins d'individus que la dimension de l'espace des données. De plus, le même phénomène que pour les régressions locales se produit si il y a trop peu d'individus pour la construction du sous espace et l'influence du bruit dans le choix des directions augmente.

3.3 Validation des approches

Dans cette partie nous mettrons en œuvre les différentes approches présentées plus haut. Nous les évaluerons sur des exemples, nous discuterons de leurs points sensibles et présenterons des résultats sur différents types de données.

3.3.1 Cas d'un problème linéaire

Prenons le cas dans lequel la variable à expliquer y est linéairement liée au vecteur de données x associé. On peut donc écrire la relation $X = y^T \beta + \epsilon$, où x est de dimension 4 et X contient 10^3 individus. Le vecteur y suit une distribution gaussienne $y \sim N(0, 1)$. Le vecteur des coefficients est $\beta^T = [0.7, 0.6, 0, -0.7]$ et le bruit $\epsilon \sim N(0, 0.1 \times I_4)$. I_4 est la matrice identité de dimension 4.

Le problème étant posé, on peut regarder l'expression de y en fonction des différentes composantes de x dans l'espace des données brutes. Ainsi, on obtient les nuages de points représentés figure 3.6.

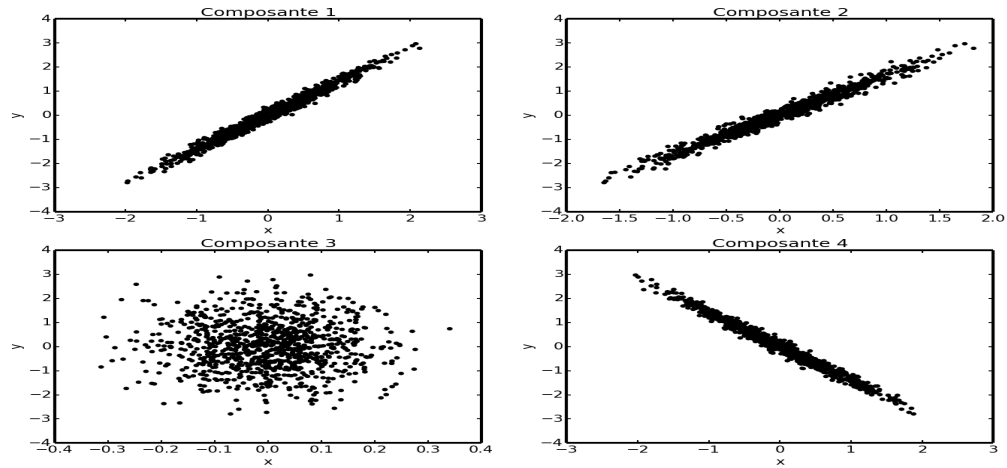


FIGURE 3.6 – Expression de la variable y en fonction des différentes composantes de x (cas de l'exemple linéaire). On voit que les composantes 1, 2 et 4 sont bien corrélées avec y et la composante 3 en revanche semble n'avoir aucun lien avec y . En effet pour une valeur de x donnée on ne peut pas déduire une valeur de y à partir de la composante 3.

On peut voir sur la figure 3.6 que les composantes 1, 2 et 4 présentent un lien linéaire avec y . Il est donc possible d'avoir une estimation de y connaissant les composantes 1, 2 ou 4 de x . L'ACP devrait nous permettre d'utiliser au mieux l'information contenue dans

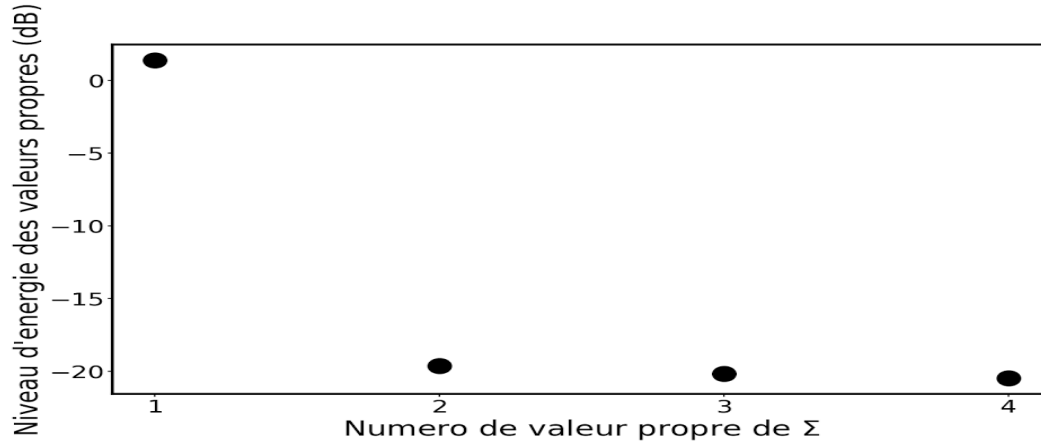


FIGURE 3.7 – Décroissance des valeurs propres de Σ pour l'exemple linéaire du paragraphe 3.3.1.

chacune des composantes.

En décomposant la matrice Σ on obtient les valeurs propres présentées figure 3.7.

On peut remarquer, sur la figure 3.7, la “cassure”, évoquée dans la section 3.1.1, dans la décroissance des valeurs propres. Cela permet de déduire qu'un espace de dimension 1 est suffisant pour résumer les données au sens de l'ACP.

Si l'on projette les données $\mathbf{X}$ sur l'espace construit par les vecteurs propres associés aux valeurs propres de la figure 3.7, on peut appliquer la même représentation que sur la figure 3.6, mais avec les données projetées sur les axes donnés par l'analyse en composante principale à la place des composantes de l'espace d'origine.

La figure 3.8 montre une relation plus marquée entre y et le comportement de x projeté sur la première composante donnée par l'ACP. On observe de manière qualitative, pour cet exemple, que sur la représentation “PCA 1” de la figure 3.8 le nuage de points est plus proche d'une droite marquant une relation linéaire entre x et y que cela pouvait l'être dans la figure 3.6. Les trois autres composantes de l'ACP ne montrent, quant à elles, aucune corrélation entre l'expression de x et de y .

Si l'on augmente le niveau de bruit, le lien entre y et x devient moins marqué comme le montre la figure 3.9. Dans ce cas, la puissance du bruit a été augmentée de 30 dB.

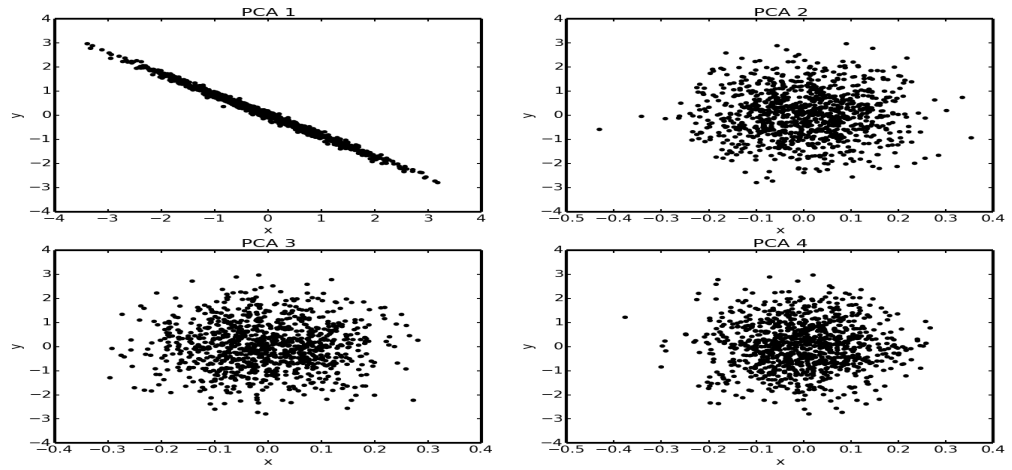


FIGURE 3.8 – Expression de y en fonction des valeurs des individus projetés sur les différentes composantes de l'ACP (cas de l'exemple linéaire du paragraphe 3.3.1)

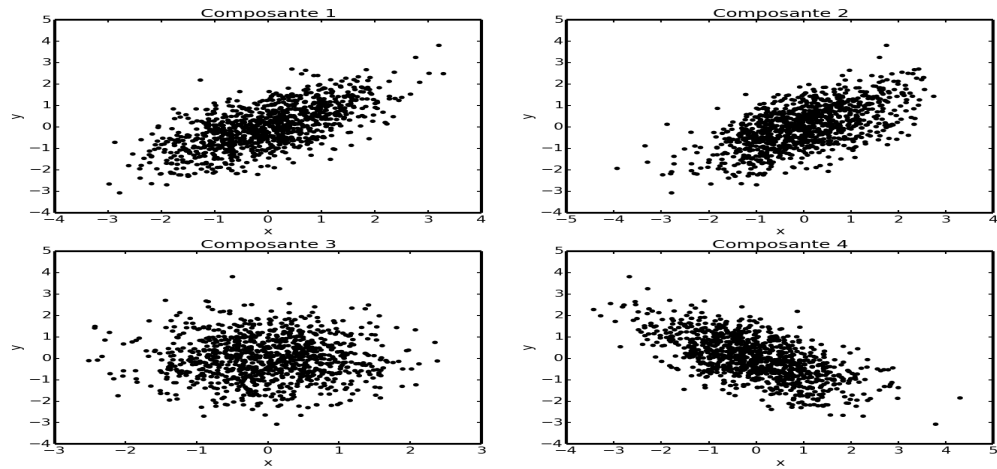


FIGURE 3.9 – Expression de y en fonction des différentes composantes des données dans l'espace d'origine pour l'exemple linéaire décrit au paragraphe 3.3.1 avec un bruit 30 dB plus puissant que dans la figure 3.6.

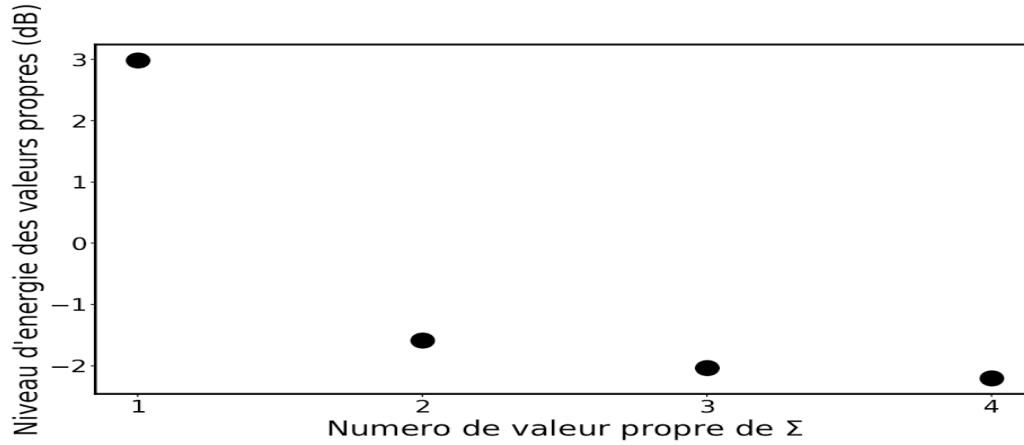


FIGURE 3.10 – Décroissance des valeurs propres de Σ pour l'exemple linéaire présenté au paragraphe 3.3.1 avec un bruit 30 dB plus puissant que dans la figure 3.6.

Les valeurs propres faibles voient leur valeur augmenter avec le bruit (figure 3.10) et bien que la cassure soit toujours nette, la seconde valeur propre est quatre fois plus faible que la première alors que le rapport entre les deux était de 120 lorsque l'on avait un faible niveau de bruit (cf. figure 3.7).

L'ACP, moins efficace que précédemment, parvient tout de même à extraire une relation entre y et x sur sa première composante (figure 3.11).

Nous pouvons à présent utiliser les différentes méthodes d'estimation présentées plus haut sur les données dans l'espace de départ (avant la projection) et dans l'espace donné par l'ACP afin d'analyser leurs comportements. Nous pouvons aussi valider de manière quantitative le traitement par les diverses méthodes de régression basées sur l'ACP que nous avons présentées. Par la suite, nous utiliserons, 10^3 individus comme base de référence à l'aide de laquelle nous estimerons la valeur de y pour les 10^2 de la base de test. ces bases sont générées de la même manière qu'au début de cette partie.

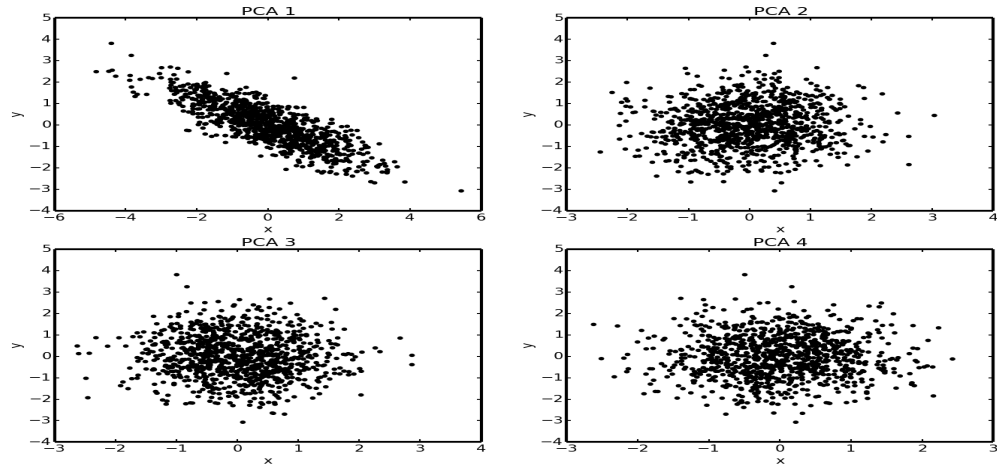


FIGURE 3.11 – Expression de y en fonction des valeurs des individus projetés sur les différentes composantes de l'ACP (cas de l'exemple linéaire avec un fort niveau de bruit).

Régressions locales

Avec un faible niveau de bruit ($\sigma_\epsilon = 0.1$), on obtient, en appliquant les régressions locales, les résultats présentés table 3.1.

Nombre de composantes	1	2	3	Pas d'ACP
Erreur	0.064	0.066	0.070	0.070

TABLE 3.1 – Table des erreurs moyennes en valeur absolue pour l'estimation de y suivant le nombre de composantes retenues en employant les régressions locales (exemple linéaire) $\sigma_\epsilon = 0.1$. La colonne "pas d'ACP" montre les résultats obtenus en appliquant la régression directement sur les données dans l'espace de départ.

La table 3.1 montre que l'utilisation de l'ACP permet, associée à la méthode de régression locale, une réduction d'environ 10% de l'erreur d'estimation de y lorsque l'on considère une direction comme l'indiquait la décroissance des valeurs propres.

Pour le cas d'un fort niveau de bruit ($\sigma_\epsilon = 0.8$), les résultats sont présentés table 3.2.

Nombre de composantes	1	2	3	Pas d'ACP
Erreur	0.49	0.47	0.47	0.44

TABLE 3.2 – Table des erreurs moyennes en valeur absolue suivant le nombre de composantes retenues en employant les régressions locales (exemple linéaire). $\sigma_\epsilon = 0.8$.

La table 3.2 montre que l'utilisation de l'ACP n'améliore plus, et même dégrade, les résultats lorsque le niveau de bruit est trop élevé. En effet, l'ACP maximise la variance des données. Si celle-ci est principalement causée par le bruit, l'ACP perd son efficacité. On peut conclure que l'ACP couplée à une méthode de régression locale n'est pertinente qu'à haut rapport signal sur bruit.

k-plus proches voisins

Avec un faible niveau de bruit ($\sigma_\epsilon = 0.1$) on obtient, en appliquant les k-plus proches voisins, les résultats présentés table 3.3. Le nombre optimal de voisins⁸ est déterminé à l'aide d'un protocole de validation croisée.

Nombre de composantes	1	2	3	Pas d'ACD
Erreur	0.072	0.073	0.077	0.078

TABLE 3.3 – Table des erreurs moyennes en valeur absolue suivant le nombre de composantes retenues avec la méthode des k-PPV (dans cet exemple 8 a été la valeur optimale de voisins retenue) (exemple linéaire). $\sigma_\epsilon = 0.1$

L'ACP, couplée avec une méthode des k-PPV, permet une réduction de l'erreur, bien qu'elle soit moindre par rapport au cas présenté table 3.1. Sur un exemple linéaire l'emploi des k-PPV n'est pas pertinent. En effet l'intérêt de l'utilisation de cette méthode est sa capacité à gérer les non-linéarités. Lorsque les données sont linéaires les moindres carrés permettent de prendre en compte plus de données.

8. Un nombre fixe de voisins sera déterminé pour tous les individus.

Conclusion sur l'exemple linéaire

On a pu voir dans cet exemple que l'efficacité de l'analyse en composante principale dépend du niveau de bruit. Un niveau de bruit trop élevé met en défaut la méthode (figure 3.11). Dans le cas d'un lien linéaire entre les données x et la variable à expliquer (paramètre) y , les régressions locales au sens des moindres carrés sont plus efficaces que la méthode des k-PPV. Cette efficacité est liée au fait qu'une régression linéaire au sens des moindres carrés, est optimale dans le cas d'un lien globalement linéaire entre données et variable à expliquer. On peut même dans ce cas se passer de contrainte de localité sur les régresseurs.

3.3.2 Cas d'un problème non-linéaire

Penchons-nous à présent sur un cas où les données x évoluent de façon non-linéaires par rapport à y . Par exemple, nous pouvons exprimer des données d'étude ainsi :

$$x_1 = y^3 + \epsilon_1 \quad (3.3)$$

$$x_2 = \sin(3y) + \epsilon_2 \quad (3.4)$$

$$x_3 = \sqrt{3} + \epsilon_3 \quad (3.5)$$

$$x_4 = \epsilon_4 \quad (3.6)$$

où les x_i et les ϵ_i sont respectivement les composantes des données et du bruit, associées à l'individu i . Le paramètre y est distribué uniformément sur l'intervalle $[0,1]$, et les composantes du bruit suivent une distribution gaussienne d'écart type $\sigma_\epsilon = 0.03$.

La figure 3.12 montre assez clairement les fonctions non-linéaires liant les différentes composantes des données à la variable à expliquer. De même que précédemment, nous nous intéressons à la décroissance des valeurs propres de Σ pour déterminer la taille du sous-espace optimal.

Contrairement à la figure 3.7, la décroissance des valeurs propres, figure 3.13, montre assez clairement que deux composantes sont très informatives au sens de la maximisation de la variance.

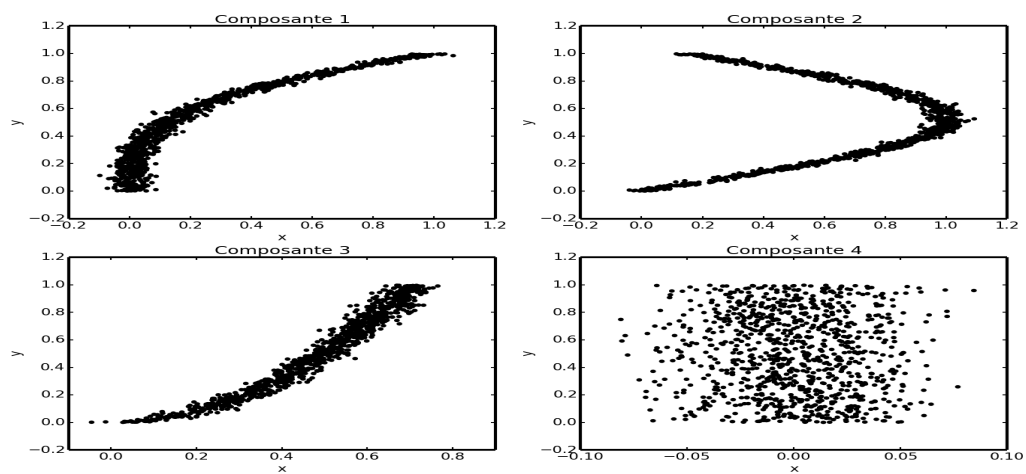


FIGURE 3.12 – Expression de y en fonction des différentes composantes des données dans l'espace d'origine pour l'exemple non-linéaire

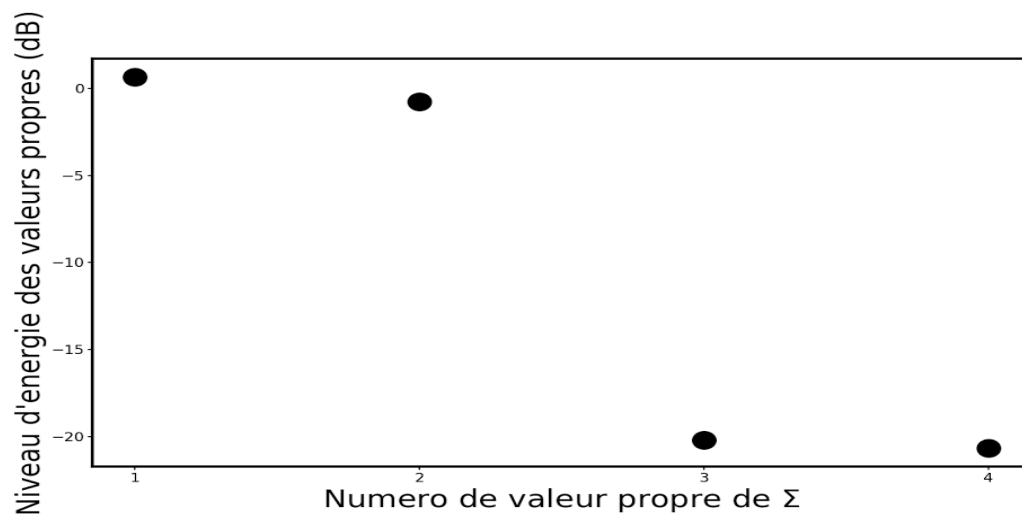


FIGURE 3.13 – Décroissance des valeurs propres de Σ pour l'exemple non-linéaire

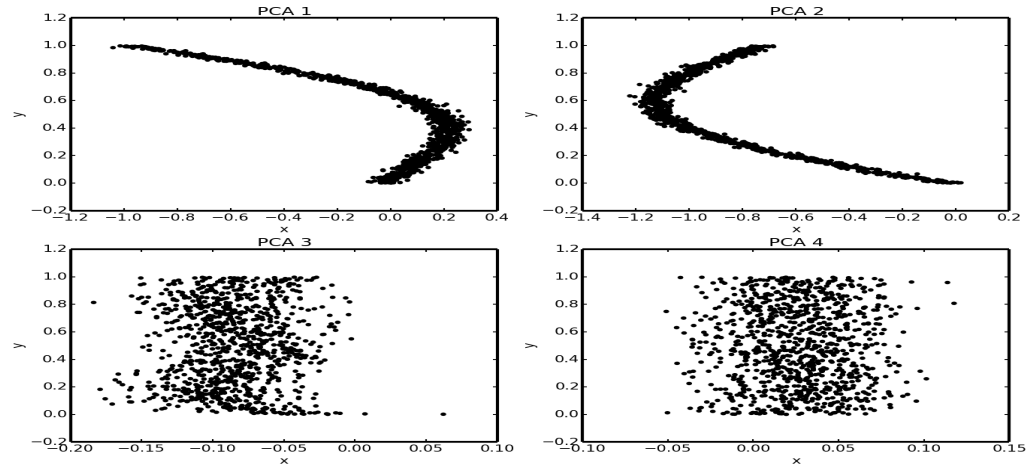


FIGURE 3.14 – Expression de y en fonction des valeurs des individus projetés sur les différentes composantes de l'ACP (cas de l'exemple non-linéaire).

Cela se confirme (cf. figure 3.14), lorsque l'on met en évidence l'évolution des valeurs de y en fonction des valeurs prises par les individus projetés sur les deux premiers axes de l'ACP : les composantes 1 et 2 montrent des formes définies tandis que la répartition suivant les composantes 3 et 4 est aléatoire.

On peut prévoir, de l'analyse des figures 3.13 et 3.14, que 2 sera la dimension optimale du sous-espace pour l'estimation de y . On peut observer sur la figure 3.14 un exemple mettant en avant l'intérêt de sélectionner les directions en fonction de l'individu traité. Par exemple, lorsqu'un individu se trouve projeté à la fois sur des valeurs faibles de PCA 1 (x autour de -1), en haut à gauche de la figure 3.14, et sur des valeurs faibles de PCA 2 (x autour de -0.8), en haut à droite de la figure 3.14, alors l'emploi des deux composantes n'est pas la meilleure approche possible. PCA 1 donne une très bonne estimation de y mais PCA 2 présente deux modes possibles. La considération du second axe, dans ce cas, n'apportera que confusion dans l'estimation de y .

Régressions locales

De même que précédemment, on applique sur l'exemple du paragraphe 3.3.2, les régressions locales. Les résultats sont présentés table 3.4.

Nombre de composantes	1	2	3	Pas d'ACP
Erreur	0.16	0.048	0.054	0.069

TABLE 3.4 – Table des erreurs moyennes en valeur absolue suivant le nombre de composantes retenues. $\sigma_\epsilon = 0.1$

La table 3.4, confirme le pronostic suivant lequel deux composantes sont optimales pour l'estimation de y . L'ACP, combinée à la méthode d'estimation basée sur les régressions locales permet de réduire de 20% l'erreur moyenne d'estimation sur cet exemple non-linéaire. Le voisinage optimal est dans ce cas déterminé par validation croisée. Celui-ci est donc le même quelle que soit la position dans l'espace. Cette approche est simple de mise en œuvre mais est localement sous-optimale.

k-plus proches voisins

De la même manière, on applique la méthode des k-PPV pour laquelle les résultats sont présentés table 3.5.

Nombre de composantes	1	2	3	Pas d'ACP
Erreur	0.13	0.040	0.050	0.55

TABLE 3.5 – Table des erreurs moyennes en valeur absolue suivant le nombre de composantes retenues. $\sigma_\epsilon = 0.05$.

Les résultats présentés table 3.5 sont meilleurs avec les k-PPV qu'avec les régressions locales. Les k-PPV semblent par ailleurs moins sensibles à la considération de composantes supplémentaires. Les k-PPV sont une approche plus pertinente pour le traitement des données non-linéaires que les régressions locales, car cette première méthode est moins

sensibles lorsqu'on ne considère pas exactement le bon nombre de directions, et qu'ils sont aussi moins sensibles au bruit.

Sélection de directions pertinentes

Ici, la sélection des directions présentées au paragraphe 3.2.4 prend son sens. En effet, la figure 3.14 montre que les deux premières directions sont globalement pertinentes, mais dans le cas, par exemple, où un individu prendrait une valeur élevée de x à la fois sur les deux axes PCA 1 et PCA 2, il est clair que l'axe PCA 1 n'apporterait pas d'information, voire pire, il pourrait apporter du bruit. Ainsi, si l'on apporte la sélection des directions associées aux régressions locales, on peut obtenir les résultats présentés table 3.6.

Realisation	Avec sélection	Sans sélection	Direction 1	Direction 2	Directions 1+2
1	0.023	0.025	7	18	75
2	0.025	0.026	11	9	80
3	0.022	0.023	18	3	79
4	0.022	0.027	8	21	71
5	0.022	0.023	13	17	70
6	0.020	0.022	23	10	67
7	0.019	0.018	13	18	69
8	0.023	0.023	14	12	74
9	0.021	0.022	20	7	73
10	0.022	0.022	14	7	79
moyenne	0.0219	0.0231	14.1	12.2	73.7

TABLE 3.6 – Comparaison des erreurs obtenues avec ou sans sélection des directions pertinentes sur 10 réalisations de l'exemple du paragraphe 3.3.2 avec la méthode utilisant les régressions locales. Les trois colonnes de droite représentent le nombre d'individus par essai dans chaque configuration

Sur cet exemple simple, on peut remarquer une légère amélioration en moyenne de l'erreur d'estimation lorsque l'on sélectionne le nombre optimal de directions indépendamment pour chaque individu. Seule la réalisation 7 montre de meilleurs résultats sans sélection. On remarque aussi qu'en moyenne, pour 26% des individus, il est préférable de construire le régresseur en vue d'une estimation sur seulement une composante plutôt que deux. L'apport global de la sélection de directions sur cet exemple n'est pas flagrant car on doit choisir entre l'espace 1D et 2D, cela donne sensiblement les mêmes résultats. Des exemples dans des espaces plus grands présenteront un choix plus varié de directions et cela permettra d'évaluer l'utilité de cette sélection.

Realisation	Avec sélection	Sans sélection	Direction 1	Direction 2	Directions 1+2
1	0.023	0.023	11	10	79
2	0.025	0.026	17	14	69
3	0.022	0.023	19	9	72
4	0.022	0.027	10	10	80
5	0.022	0.023	16	10	74
6	0.020	0.022	20	13	67
7	0.019	0.018	20	6	74
8	0.023	0.023	16	6	78
9	0.021	0.022	11	18	71
10	0.022	0.022	10	14	76
moyenne	0.0230	0.0232	15	11	74

TABLE 3.7 – Comparaison des erreurs obtenues avec ou sans sélection des directions pertinentes sur 10 réalisations avec la méthode utilisant les k-plus proches voisins. Les trois colonnes de droite représentent le nombre d'individus par essai

La table 3.7 présente les résultats sur le même exemple avec l'emploi des k-PPV plutôt que les régressions locales. On peut observer que, dans ce cas, la sélection des directions semble apporter un gain beaucoup moins important, et surtout moins systématique. Néan-

moins, en moyenne, on constate une diminution même faible de l'erreur d'estimation.

Conclusion sur l'exemple non-linéaire

Sur cet exemple, on peut constater à la fois la nécessité de sélectionner plusieurs composantes de l'espace engendré par l'ACP et l'intérêt de la sélection des composantes en fonction du voisinage de l'individu traité.

3.3.3 Cas d'un problème complexe

Les exemples précédents avaient l'avantage d'être dans des espaces de faible dimension, donc aisément représentables. Cependant, les données que nous souhaitons traiter sont dans des espaces de grande dimension (de l'ordre de 10^4 par exemple). Nous allons donc à présent étudier ces méthodes sur un exemple plus complexe : avec des données dans un espace de dimension 200.

Nous utiliserons des pseudo-raies spectrales uniques simulées à l'aide de gaussiennes, toutes centrées sur la même longueur d'onde, mais avec des largeurs et dépressions différentes. Les largeurs et les dépressions de ces gaussiennes seront les paramètres que nous estimerons à partir des vecteurs de données. L'équation 3.7 ci-dessous, montre que le paramètre de dépression s'exprime linéairement dans les données, ce qui n'est pas le cas du paramètre de largeur. Les vecteurs de données x associés à chaque individu suivent l'expression :

$$x(\lambda; h, w) = 1 - h \times e^{-\frac{1}{2} \left(\frac{\lambda - \lambda_0}{w} \right)^2} + \epsilon. \quad (3.7)$$

où λ représente l'espace (spectral au sens physique du terme) d'expression des données, λ_0 est la position de la raie commune à tous les individus, h est le paramètre de dépression et w le paramètre de largeur ; ϵ est un bruit additif gaussien.

Nous pouvons appliquer une analyse en composante principale et observer l'expression des données en fonction de chacun des paramètres de dépression et de largeur.

Pour différentes valeurs du rapport signal sur bruit (RSB⁹), on obtient les résultats présentés table 3.8.

9. $RSB = 10 \log_{10}(P_s/P_b)$ où P_s est la puissance du signal et P_b la puissance du bruit.

	k-PPV	k-PPV + sélection	Régression locale	Régression + sélection
RSB = 25dB				
Dépressions	0.0082	0.0088	0.020	0.019
Largeurs	0.020	0.022	0.044	0.037
RSB = 11dB				
Dépressions	0.014	0.017	0.037	0.038
Largeurs	0.049	0.049	0.088	0.070
RSB = 6dB				
Dépressions	0.026	0.028	0.044	0.037
Largeurs	0.080	0.087	0.12	0.10

TABLE 3.8 – Erreurs moyennes en valeurs absolues obtenues pour l’estimation des dépressions et largeurs des gaussiennes pour différents niveaux de bruit et avec les différentes méthodes présentées.

La table 3.8 montre que la sélection des directions n’apporte rien (voire dégrade les résultats) à une approche par k-PPV. En effet, le voisinage défini permet à la méthode par k-PPV d’avoir une bonne prise en charge des non-linéarités de cet exemple. Cette sélection permet de réduire légèrement l’erreur d’estimation dans le cas d’une approche par régressions locales. Cependant, comme dans le cas de l’exemple du paragraphe 3.3.2, l’approche par régressions locales ne semble pas présenter d’intérêt pour cet exemple.

3.4 Application aux données astrophysiques

Nous pouvons désormais appliquer les méthodes basées sur l’ACP, qui viennent d’être présentées, à des données propres au domaine d’application. Nous utiliserons alors une base de données de spectres synthétiques (cf. figure 1.4), pour lesquels les individus sont représentés dans un espace de dimension 8000 correspondant aux 8000 longueurs d’ondes considérées. Nous chercherons à déterminer les valeurs des paramètres fondamentaux qui

sont la température effective T_{eff} , la gravité de surface $\log(g)$, la métallicité $[Fe/H]$ ainsi que les valeurs de la vitesse de rotation projetée $v \sin(i)$ ¹⁰.

La décroissance des valeurs propres ne permet pas de déterminer combien de directions il est pertinent de garder. En effet, on n’observe pas de cassure dans la décroissance des valeurs propres. Il est possible, en revanche, de placer un seuil à 99% de la conservation d’inertie et cela revient à conserver les 12 premières valeurs propres (cf. figure 3.15).

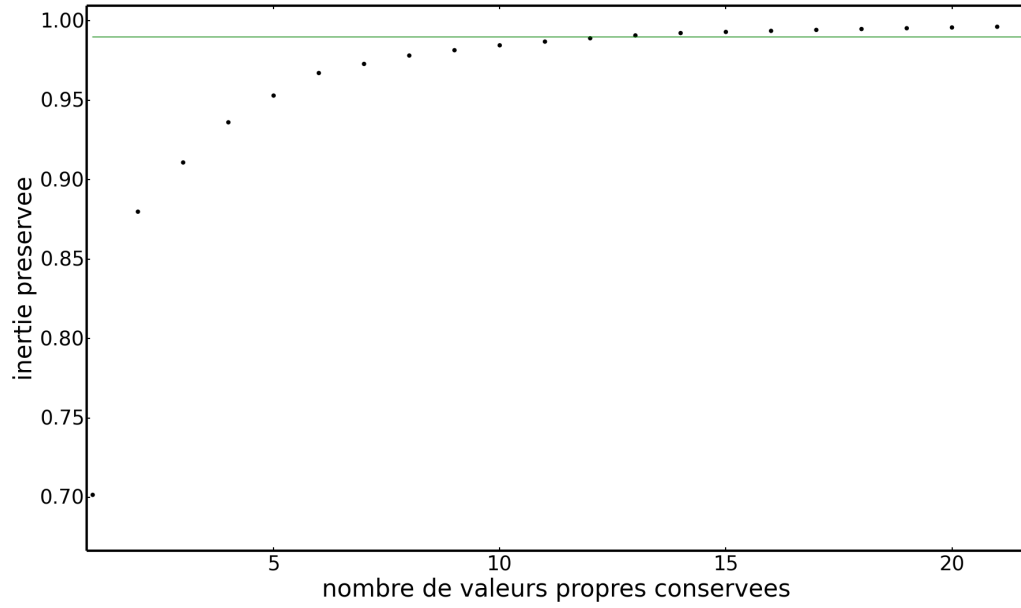


FIGURE 3.15 – Préservation de l’inertie en fonction du nombre de valeurs propres conservées pour la base de données de spectres synthétiques.

On peut alors, pour différents niveaux de bruit, comparer les valeurs des erreurs pour les différents paramètres (cf. figure 3.16).

Les résultats présentés figure 3.16 confirment ce qui avait été observé précédemment : les

10. Les spectres couvrent le domaine spectral de 390 nm à 680 nm avec des valeurs de température allant de 4000 K à 7500 K, des valeurs de gravité de surface allant de 3 à 5 dex des valeurs de métallicité allant de -1 à 1.5 dex et des valeurs de vitesse de rotation projetée allant de 0 à 150 km/s.

régressions locales montrent leur intérêt à haut rapport signal sur bruit (dans cet exemple utilisant les spectres synthétiques issus du modèle d’atmosphère Atlas, le bruit est gaussien et indépendant, ce qui n’est pas le cas des données réelles). Une exception est faite pour le paramètre de métallicité ($[Fe/H]$) qui est presque linéaire par rapport aux données (il se traduit directement de la “richesse spectrale”, c’est-à-dire la multiplicité des petites raies).

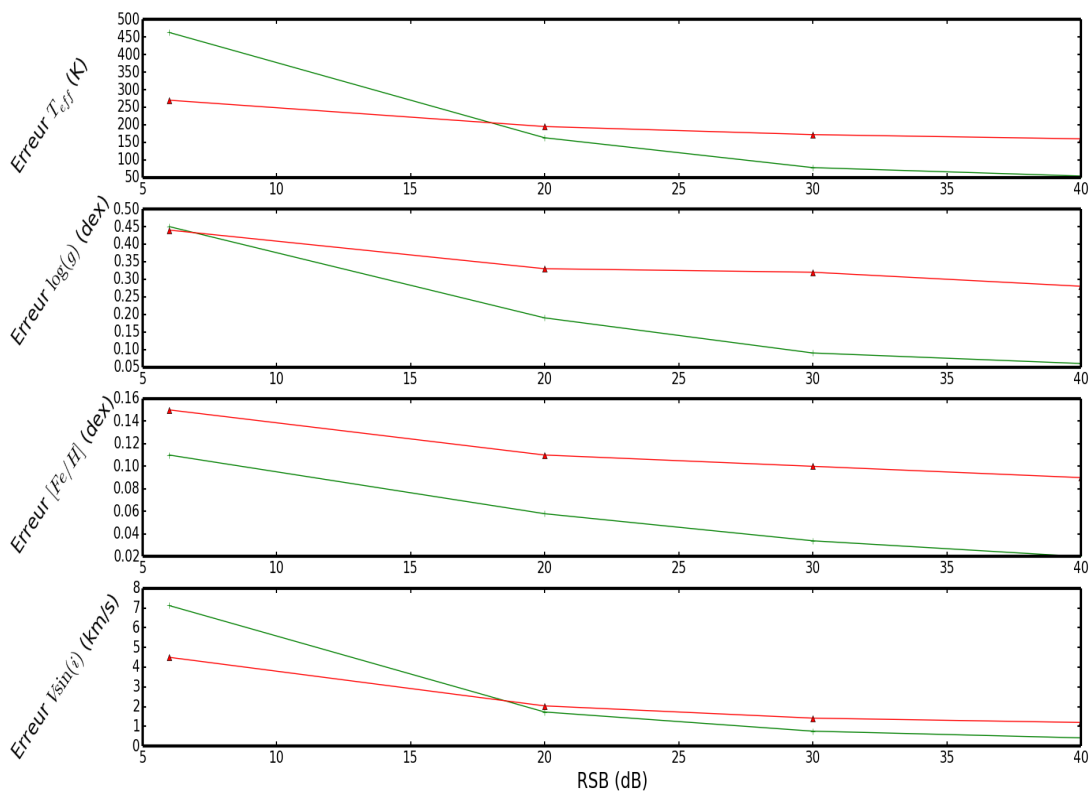


FIGURE 3.16 – Erreurs d’estimation pour les différents paramètres en fonction du RSB avec les k-PPV (rouge) et les régressions locales (vert). Les k-PPV donnent de meilleurs résultats pour la température et la métallicité à faible RSB.

3.5 Sélection de directions sur données observées

Dans ce dernier paragraphe, nous travaillerons sur les données utilisées par Paletou et al. (2015a). Ces données sont constituées de deux bases de spectres observés. Une de ces bases de spectres, la base "Elodie" (car les spectres de cette base sont tous issus de cet instrument), sert de base de données de référence. Elle est constituée de 905 spectres de types F, G et K¹¹. La seconde base, la base de test que l'on appellera "S4N", est une base contenant 104 spectres de même types que pour la base "Elodie" (Allende Prieto et al., 2004). Dans cette partie nous exprimerons les erreurs en termes de biais et d'écart type pour faciliter la comparaison avec les résultats de l'article de référence (Paletou et al., 2015a).

3.5.1 Température effective (T_{eff})

De la même manière que dans le paragraphe 3.3, sont représentés figure 3.17 les individus de la base de référence avec en abscisses la valeur de leur projection sur une des composantes de l'ACP, et en ordonnée les valeurs prises par la température effective T_{eff} .

La première composante de l'ACP représentée figure 3.17 présente une forte corrélation avec les valeurs prises par la température, et cela laisse à penser que celle-ci est toujours pertinente pour l'estimation des valeurs de la température effective. En revanche, la seconde direction ne paraît que peu pertinente pour l'estimation de valeurs de températures effectives qui seraient supérieures à 5000 K (les deux tiers les plus chauds des individus donnent tous une valeur de $x^T v_2$ autour de 0). Pour les autres directions, il est difficile de juger *a priori* leurs pertinences, car les formes des nuages de points sont moins marquées. En appliquant la sélection des directions comme présentée au paragraphe 3.2.4, on détermine (figure 3.18) la dimension des sous-espaces les plus pertinents pour chaque individu.

La figure 3.18 montre que les sous-espaces pertinents sont de dimension 11 ou 12 pour près de 60% des individus. Cela montre que, dans le cas de la température, le sous-espace choisi dans l'article de Paletou et al. (2015a), de dimension 12, est bien le plus pertinent.

11. Le chapitre 1 décrit les types spectraux.

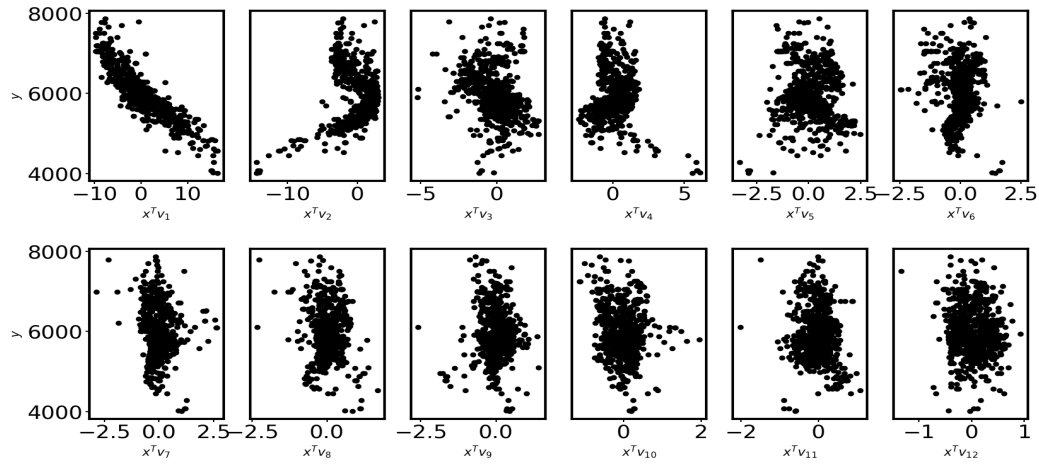


FIGURE 3.17 – Projection des individus de la base Elodie sur les 12 premières composantes de l'ACP. En ordonnées les valeurs correspondantes de T_{eff} .

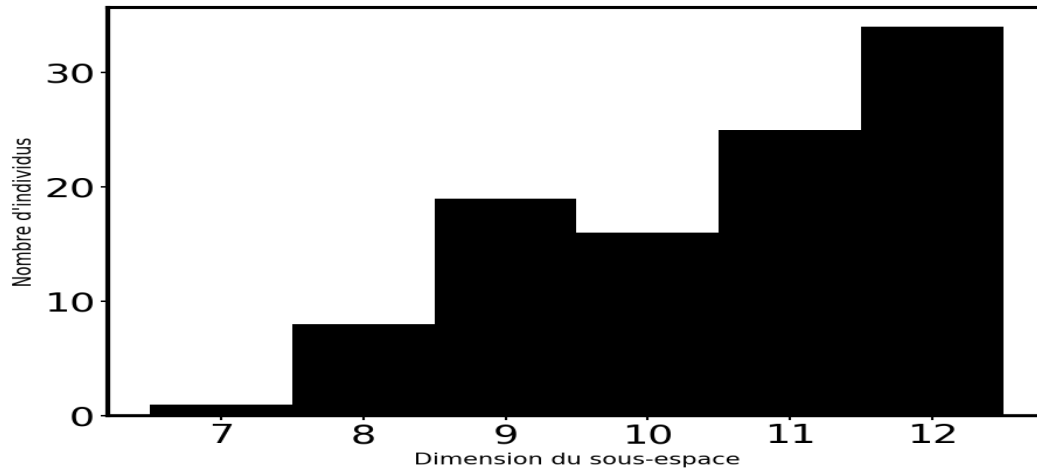


FIGURE 3.18 – Répartition des dimensions de sous-espace choisies pour la projection des individus lors du traitement de T_{eff} .

	Sans sélection	Avec sélection
Biais	-56	-53
Ecart-type	100	115

TABLE 3.9 – Biais et écart type de l’erreur d’estimation de T_{eff} (K), pour l’échantillon des spectres du S4N avec et sans la sélection des directions. La première colonne présente les résultats de l’article de référence (Paletou et al., 2015a).

Ainsi, nous pouvons comparer les résultats obtenus en terme de biais et de variance avec l’application de la méthode proposée dans Paletou et al. (2015a) table 3.9.

Dans le cas de la température, la sélection des directions de l’ACP n’est pas probante en moyenne. Le biais de l’estimation est réduit, mais l’écart-type augmente.

3.5.2 Gravité de surface ($\log(g)$)

L’expression de la gravité $\log(g)$ sur les directions de l’ACP est présentée figure 3.19.

Sur la figure 3.19, il est beaucoup plus difficile que sur la figure 3.17 de voir une corrélation évidente entre la projection des individus sur l’une des composantes de l’ACP et la valeur prise par la gravité de surface. Cependant la méthode de sélection de direction doit composer, à partir d’une combinaison optimale des directions, le sous espace pour chaque individu.

La figure 3.20 montre que, par rapport à la figure 3.18, beaucoup d’individus (environ 75 sur 104) sont mieux représentés sur des sous-espaces de dimension inférieure à 10. On peut en déduire que le sous-espace de dimension 12 choisi originellement pour l’ACP n’est pas aussi pertinent pour la gravité de surface qu’il ne l’était pour la température effective.

La table 3.10 montre une réduction de 5% de l’erreur en terme d’écart type. Cela montre que l’évolution du paramètre peut, statistiquement, être mieux caractérisée sur un espace de dimension plus faible déterminé comme localement optimal. De manière pragmatique, la sélection des directions ne produit pas une réduction flagrante de l’erreur, mais tout de même, cela nous permet de vérifier statistiquement sur ces données l’hypothèse de départ qui est que le sous-espace le plus pertinent n’est pas le même pour tous les individus.

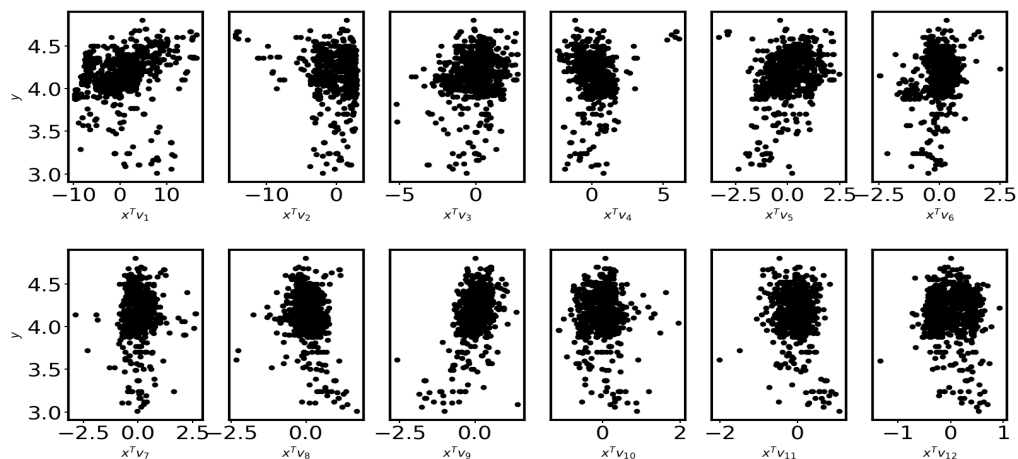


FIGURE 3.19 – Projection des individus de la base Elodie sur les 12 premières composantes de l’ACP. En ordonnées les valeurs correspondantes de $\log(g)$.

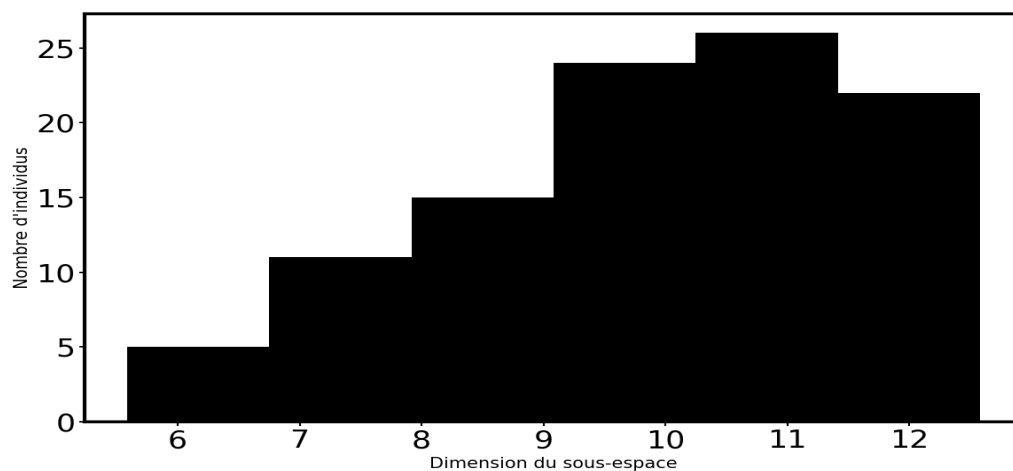


FIGURE 3.20 – Répartition des dimensions de sous-espace choisies pour la projection des individus lors du traitement de $\log(g)$.

	Sans sélection	Avec sélection
Biais	0.19	0.19
Ecart-type	0.18	0.17

TABLE 3.10 – Biais et écart type de l’erreur d’estimation de $\log(g)$ (dex), pour l’échantillon des spectres du S4N avec et sans la sélection des directions. La première colonne présente les résultats de l’article de Paletou et al. (2015a).

	Sans sélection	Avec sélection
Biais	0.01	0.006
Ecart-type	0.11	0.12

TABLE 3.11 – Biais et écart type de l’erreur d’estimation de $[Fe/H]$ (dex), pour l’échantillon des spectres du S4N avec et sans la sélection des directions. La première colonne présente les résultats de l’article de Paletou et al. (2015a).

3.5.3 Métallicité ($[Fe/H]$)

L’expression de la métallicité $[Fe/H]$ sur les directions de l’ACP sest présentée figure 3.21.

Dans ce cas, il est intéressant de noter que les premières composantes (figure 3.21) ne semblent pas être les plus pertinentes. En effet, les composantes 3 et 4 présentent une corrélation¹² avec la métallicité qui semble plus importante.

Comme pour la température effective, on peut constater pour la métallicité que près de 60% des individus se voient mieux représentés dans un espace de dimension 11 ou 12 (cf. figure 3.22).

Pour la métallicité, on a, comme dans le cas de la température, une légère dégradation des résultats en termes d’écart type (table 3.11). Ce résultat est cohérent avec les résultats de la figure 3.22. Bien que les composantes 3 et 4 semblent plus pertinentes, il reste

12. La corrélation apparente est donnée par l’épaisseur du nuage de points tout en tenant compte de sa pente. Plus le nuage de points est étroit avec une faible pente plus l’estimation de y sachant x est précise.

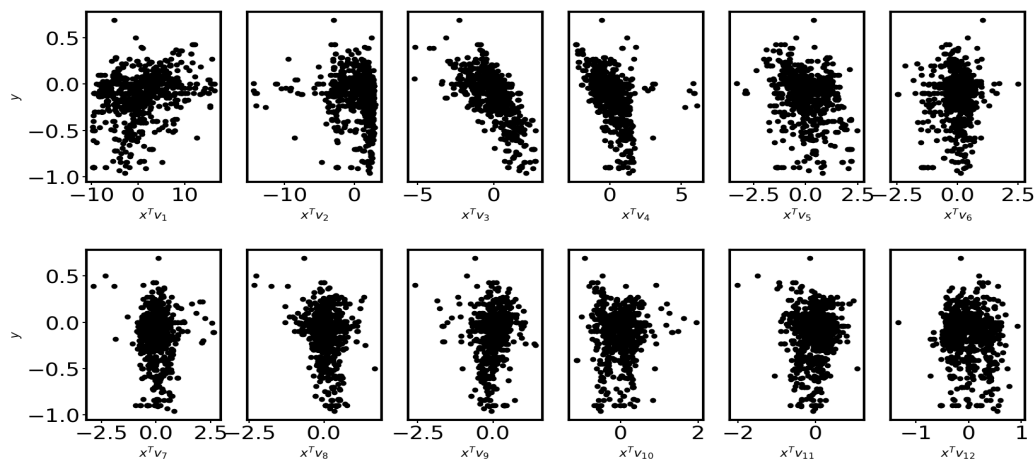


FIGURE 3.21 – Projection des individus de la base Elodie sur les 12 premières composantes de l’ACP. En ordonnées les valeurs correspondantes de $[Fe/H]$.

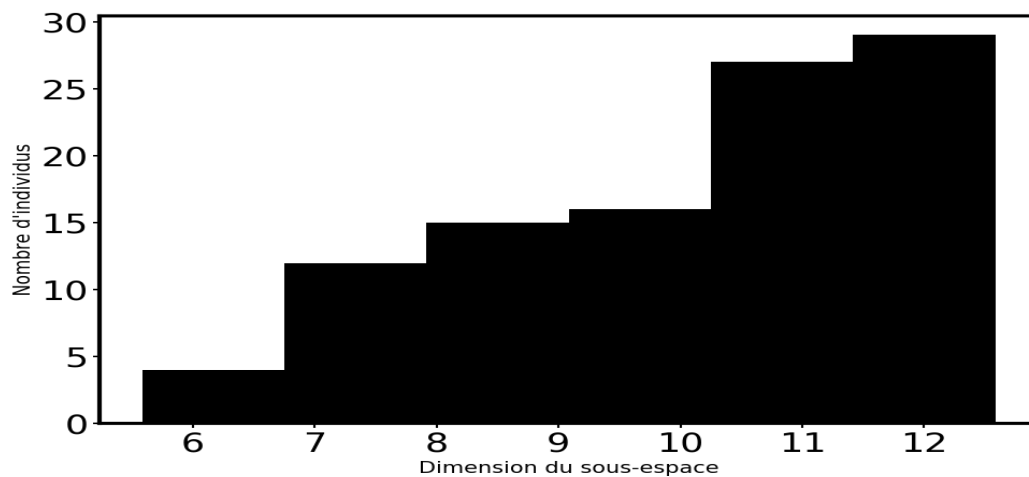


FIGURE 3.22 – Répartition des dimensions de sous-espace choisies pour la projection des individus lors du traitement de $[Fe/H]$.

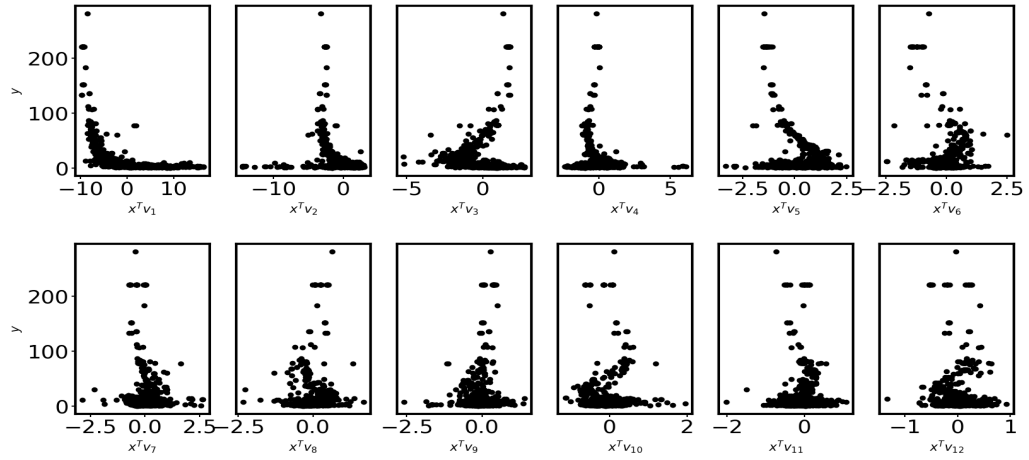


FIGURE 3.23 – Projection des individus de la base Elodie sur les 12 premières composantes de l’ACP. En ordonnées les valeurs correspondantes de $v \sin(i)$

préférable de conserver ici, presque tout le temps, toutes les composantes. Dans le cas de l’utilisation de l’ACP comme méthode de projection, il n’est pas pertinent de sélectionner les directions pour l’estimation de la métallicité.

3.5.4 Vitesse de rotation projetée ($v \sin(i)$)

La figure 3.23 présente la projection des individus sur les 12 premières composantes de l’ACP avec en ordonnées la valeur de la vitesse de rotation projetée sur l’axe de visée $v \sin(i)$ correspondant.

Il est très difficile de voir (figure 3.23) quelles directions sont les plus informatives quant à la valeur de la vitesse de rotation projetée. La composition des différentes composantes est informative, mais il n’est pas possible (figure 3.23) de déterminer visuellement une ou des directions plus pertinentes que d’autres.

La diminution du biais (cf. table 3.12) ne représente qu’une faible amélioration (car l’erreur en termes d’écart type est d’un ordre de grandeur supérieur). Comme pour les précédents, on peut représenter les individus sur des sous-espaces de dimension variable

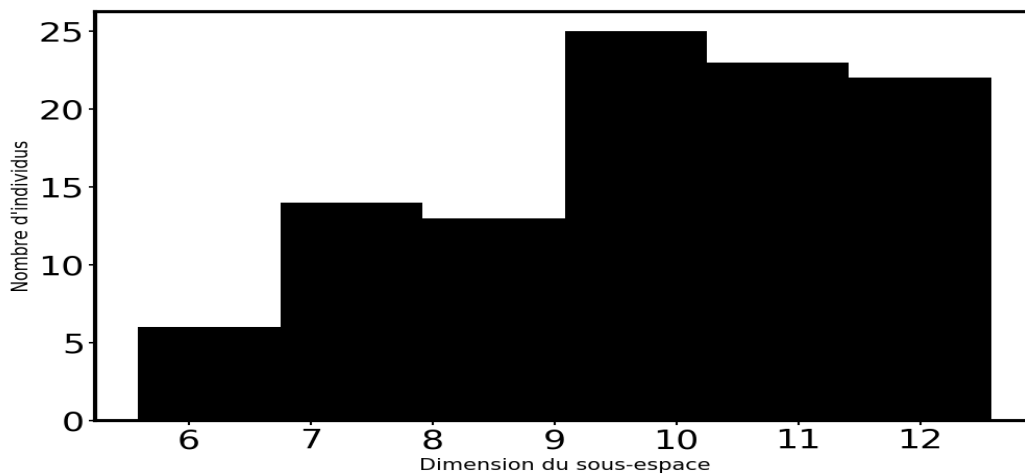


FIGURE 3.24 – Répartition des dimensions de sous-espace choisies pour la projection des individus lors du traitement de $v \sin(i)$

	Sans sélection	Avec sélection
Biais	-0.46	-0.26
Ecart-type	1.69	1.69

TABLE 3.12 – Biais et écart type de l’erreur d’estimation de $v \sin(i)$ (km/s), pour l’échantillon des spectres du S4N avec et sans la sélection des directions. La première colonne présente les résultats de l’article de Paletou et al. (2015a).

(cf. figure 3.12) sans dégradation des résultats, mais il n’y a un gain qu’au niveau du biais.

3.6 Conclusion

L’approche par sélection de direction ne permet pas de réduire l’erreur dans le cadre de l’utilisation de l’ACP. Cela permet de réduire, sans perte d’information, l’expression de certains individus à un sous-espace de dimension inférieure (sans perte d’information, statistiquement, car pas d’augmentation de l’erreur d’estimation). L’hypothèse, avancée au paragraphe 3.2.4, selon laquelle le sous-espace optimal n’est pas le même pour tous les individus, et que l’on peut l’approcher par une sélection individuelle des directions est vérifiée car il n’y a pas d’augmentation statistique de l’erreur d’estimation. Cependant dans le cadre de l’ACP cela n’apporte pas d’amélioration pour le traitement des données spectroscopiques. L’article qui clos ce chapitre montre l’application de l’ACP et des k-plus proches voisins sur les données réelles sans l’ajout de la sélection des directions.

Inversion of stellar fundamental parameters from Espadons and Narval high-resolution spectra, dans le journal *Astronomy and astrophysics*.

Résumé : Cette étude concerne l'inversion de paramètres stellaires fondamentaux, à partir de spectres à haute résolution. En effet, nous souhaitons développer un outil rapide et fiable pour le traitement des spectres obtenus grâce aux spectropolarimètres Espadons et Narval. Il permet d'effectuer une réduction de dimensionnalité et définit une métrique spécifique pour la recherche des plus proches voisins entre un spectre observé et une sélection de spectres observés issus de la bibliothèque Elodie. La température effective, la gravité de surface, la métallicité totale et la vitesse de rotation projetée sont alors déduits. Quelques tests sont présentés dans cette étude, et sont faits à partir de la seule information issue de la bande spectrale centrée autour du triplet b du magnésium et sur des étoiles de type FGK et sont très prometteurs.

Inversion of stellar fundamental parameters from ESPaDOnS and Narval high-resolution spectra[★]

F. Paletou^{1,2}, T. Böhm^{1,2}, V. Watson^{1,2}, and J.-F. Trouilhet^{1,2}

¹ Université de Toulouse, UPS-Observatoire Midi-Pyrénées, IRAP, 31400 Toulouse, France
 e-mail: fpaletou@irap.omp.eu

² CNRS, Institut de Recherche en Astrophysique et Planétologie, 14 av. É. Belin, 31400 Toulouse, France

Received 4 August 2014 / Accepted 10 November 2014

ABSTRACT

The general context of this study is the inversion of stellar fundamental parameters from high-resolution Echelle spectra. We aim at developing a fast and reliable tool for the post-processing of spectra produced by ESPaDOnS and Narval spectropolarimeters. Our inversion tool relies on principal component analysis. It allows reducing dimensionality and defining a specific metric for the search of nearest neighbours between an observed spectrum and a set of observed spectra taken from the Elodie stellar library. Effective temperature, surface gravity, total metallicity, and projected rotational velocity are derived. Various tests presented in this study that were based solely on information coming from a spectral band centred on the Mg b-triplet and had spectra from FGK stars are very promising.

Key words. astronomical databases: miscellaneous – methods: data analysis – stars: fundamental parameters

1. Introduction

This study is concerned with the inversion of fundamental stellar parameters from the analysis of high-resolution Echelle spectra. We focus on data collected since 2006 with the Narval spectropolarimeter mounted on the 2 m aperture *Télescope Bernard Lyot* (TBL) located at the summit of the Pic du Midi de Bigorre (France) and on data collected since 2005 with the ESPaDOnS spectropolarimeter mounted on the 3.6 m aperture CFHT telescope (Hawaii). We investigate, in particular, the capabilities of the principal component analysis (hereafter PCA) for setting up a fast and reliable tool to invert of stellar fundamental parameters from these high-resolution spectra.

The inversion of stellar fundamental parameters for each target that was observed with both Narval and ESPaDOnS spectropolarimeters constitutes an essential step towards (i) the subsequent post-processing of the data such as extracting polarimetric signals (see e.g., Paletou 2012 and references therein), but also (ii) exploring, or “data mining”, of the full set of data accumulated over the past eight years. In Sect. 2, we briefly describe the actual content of this database¹.

The PCA has been used for stellar spectral classification since Deeming (1964). It has been in use since, and more recently for the purpose of inverting stellar fundamental parameters from the analysis of spectra of various resolutions. It is most often used together with artificial neural networks, however (see e.g., Bailer-Jones 2000; Re Fiorentin et al. 2007). The PCA is

used there to reduce the dimensionality of the spectra before working on a multi-layer perceptron, which in turn allows linking input data to stellar parameters.

Our usage of the PCA for the inversion process is strongly influenced by this routinely made in solar spectropolarimetry in the past decade after the pioneering work of Rees et al. (2000). Very briefly, the reduction of dimensionality allowed by the PCA is directly used to build a specific metric from which a nearest-neighbour(s) search is made between an observed data set and the content of a training database. The latter can be made of synthetic data generated from input parameters that properly cover the a priori range of physical parameters expected to be deduced from the observations themselves. A quite similar use of PCA was also presented to classify and estimate of redshift of galaxies by Cabanac et al. (2002). However, in this study we use observed spectra from the Elodie stellar library as our reference data set (Prugniel et al. 2007). These spectra have for instance been used in a recent study related to the determination of atmospheric parameters of FGKM stars in the Kepler field (Molenda-Žakowicz et al. 2013). Fundamental elements of our method are described in Sect. 3. Its main originality relies also on simultaneously inverting the effective temperature T_{eff} , the surface gravity $\log g$, the metallicity $[\text{Fe}/\text{H}]$, and the projected rotational velocity $v \sin i$ directly and only from a specific spectral band extracted from the full range covered by ESPaDOnS and Narval.

Compared with alternative methods such as χ^2 fitting to a library of (synthetic) spectra, as done by Munari et al. (2005) to analyse the RAVE survey, for instance, the main advantages of the PCA are that it reduces dimensionality – a critical issue when dealing with high-resolution spectra that cover a very broad bandwidth. This allows a very fast processing of the data. Another advantage is the denoising of the original data (see e.g., Bailer-Jones et al. 1998 or Paletou 2012 in another con-

[★] Based on observations obtained at the *Télescope Bernard Lyot* (TBL, Pic du Midi, France), which is operated by the Observatoire Midi-Pyrénées, Université de Toulouse, Centre National de la Recherche Scientifique (France) and the Canada-France-Hawaii Telescope (CFHT) which is operated by the National Research Council of Canada, CNRS/INSU and the University of Hawaii (USA).

¹ polarbase.irap.omp.eu

text though). The PCA also differs from the projection method Matisse, which uses specific projection vectors that are attached to each stellar parameter that is to be inverted (Recio-Blanco et al. 2006). In that frame, these vectors are derived after assuming that they are linear combinations of every item belonging to a learning database of synthetic spectra.

In this study we restrict ourselves, on purpose, to a spectral domain ranging from 500 to 540 nm, that is, around the Mg b-triplet lines. The main argument in favour of this spectral domain is that the spectral lines of this triplet are excellent surface gravity indicators (e.g., Cayrel & Cayrel 1963; Cayrel de Strobel 1969), $\log g$ being at the same time the most difficult parameter to retrieve from spectral data. It is also a spectral domain devoid of telluric lines. More recently, Gazzano et al. (2010) performed a convincing spectral analysis of Flames/Giraffe spectra working with a similar spectral domain. However, in our study we used observed spectra as the training database for our PCA-based inversion method.

Preliminary tests made by inverting spectra taken from the S⁴N survey (Allende Prieto et al. 2004) are discussed in Sect. 4. Finally, we proceed with inverting 140 spectra of FGK stars from P B for which fundamental parameters are already available from the S catalogue of Valenti & Fisher (2005).

2. Sources of data

Our reference spectra are taken from the Elodie stellar library (Prugniel et al. 2007; Prugniel & Soubiran 2001). They are publicly available² and fully documented. The wavelength coverage of Elodie spectra is about 390–680 nm. This, unfortunately, prevents us from testing in the spectral domain at the vicinity of the infrared triplet of Ca, for instance (see e.g., Munari 1999 for a detailed case concerning this spectral range). We also used the high-resolution spectra at $\mathcal{R} \sim 42\,000$.

First tests of our method were made with stellar spectra from the Spectroscopic Survey of Stars in the Solar Neighbourhood, also known as S⁴N (Allende Prieto et al. 2004). They are also publicly³ available and fully documented. The wavelength coverage is 362–1044 nm (McDonald Observatory, 2.7 m telescope) or 362–961 nm (La Silla, 1.52 m telescope), and the spectra have a resolution $\mathcal{R} \sim 50\,000$.

However, our main purpose is inverting of stellar parameters from high-resolution spectra coming from Narval and ESPaDOnS spectropolarimeters. These data are now available from the public database P B (Petit et al. 2014). Narval is a modern spectropolarimeter operating in the 380–1000 nm spectral domain, with a spectral resolution of 65 000 in its polarimetric mode. It is an improved copy, adapted to the 2 m TBL telescope, of the ESPaDOnS spectropolarimeter, which is in operations since 2004 at the 3.6 m aperture CFHT telescope.

P B is operational since 2013. It is at the present time the largest on-line archive of high-resolution polarization spectra. It hosts data that were taken at the 3.6 m CFHT telescope since 2005 and with the 2 m TBL telescope since 2006. So far, more than 180 000 independent spectra are available for more than 2 000 distinct targets all over the Hertzsprung-Russell diagram. More than 30 000 polarized spectra are also available, mostly for circular polarization. Linear polarization data are still very rare and amount to a about 2% of the available data.

At present, the P B database provides only Stokes I or V/I_c spectra calibrated in wavelength. Stokes I data are either

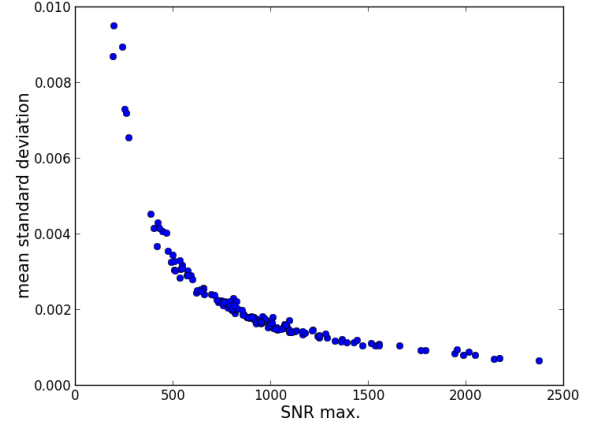


Fig. 1. Typical domain of variation of the noise level associated with the ESPaDOnS-Narval spectra that we process here. The mean standard deviation of noise per pixel for the wavelength range around the b-triplet of Mg is displayed here vs. the highest signal-to-noise ratio (SNR) of the full spectra.

normalized to the local continuum or not. We have plans to propose higher-level data however, such as pseudo-profiles resulting from line addition and/or least-squares deconvolution (see e.g., Paletou 2012), activity indexes and stellar fundamentals parameters. These last, in addition to being obviously interesting in themselves, is also indispensable for any accurate subsequent post-processing of these high-resolution spectra. These spectra also generally have high signal-to-noise ratio, as can be seen in Fig. 1. The Stokes I spectra we used result from combining four successive exposures, each of them carrying two spectra of orthogonal polarities generated by a Savart plate-type analyser. This procedure of double “beam-exchange” measurement is designed for the purpose of extracting (very) weak polarization signals (see e.g., Semel et al. 1993).

3. PCA-based inversion

Our PCA-based inversion tool is strongly inspired by magnetic and velocity field inversion tools that have been developed in the past decade to diagnose solar spectropolarimetric data (see e.g., Rees et al. 2000). Improvements of this method have recently been reported by Casini et al. (2013), for instance. Below we describe its main characteristics in the particular context of our study.

3.1. Training database

Our training database was created after the Elodie stellar spectral library (Prugniel et al. 2007; Prugniel & Soubiran 2001; see also the Vizier at CDS catalogue III/251). Stellar parameters associated with each spectrum were extracted by us from the CDS, using resources from the Python package *astroquery*⁴, in particular the components allowing queries of Vizier catalogues.

We complemented this information with a value of the projected rotational velocity for each object and spectrum of our database. To do so, we queried the catalogue of stellar rotational velocities of Głeboccki & Gnaniński (2005, also Vizier at

² See: <http://atlas.obs-hp.fr/elodie/>

³ hebe.as.utexas.edu/s4n/

⁴ astroquery.readthedocs.org

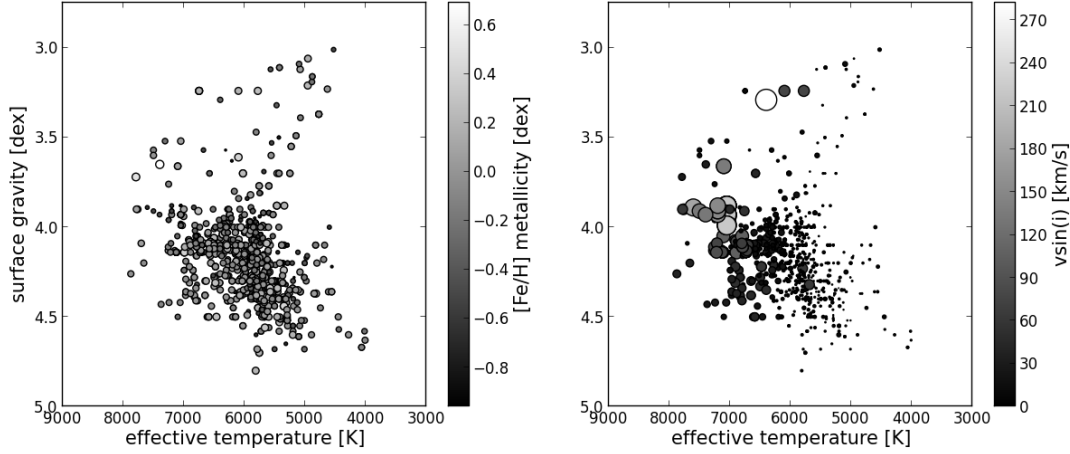


Fig. 2. Graphical summary of the coverage in stellar parameters corresponding to the content of our Elodie spectra training database (for FGK stars). Note that we adopted values of $v \sin i$ from the Vizier at CDS catalogue III/244 (Głebocki & Gnaniński 2005). Size and colour of each dot are proportional to either $[\text{Fe}/\text{H}]$ (left) or $v \sin i$ (right).

CDS catalogue III/244⁵). Moreover, we removed all objects for which we could not find a value of $v \sin i$ in this catalogue.

Since our preliminary tests concerns the inversion of the S⁴N survey as well as objects in common of P B and of the S catalogue (Valenti & Fisher 2005), we limited ourselves to spectra for which T_{eff} lies between 4000 and 8000 K, $\log g$ is greater than 3.0 dex, and $[\text{Fe}/\text{H}]$ is greater than -1.0 dex. We also had to reject a few Elodie spectra that we found to be improperly corrected for radial velocity and therefore misaligned in wavelength with respect to the other spectra. The various coverage in stellar parameters from the finally 905 selected spectra are summarized in Figs. 2.

Finally, following Muñoz Bermejo et al. (2013), we adopted the same renormalization procedure, homogeneously, for the whole set of Elodie spectra that were our training database. It is an iterative method consisting of two main steps. The first stage consists of fitting the normalized flux in the spectral bandwidth of interest by a high-order (eight) polynomial. Then we computed $D(\lambda)$ that is the difference between the initial spectra and the polynomial fit and its mean $\bar{D}$ and standard deviation σ_D . We rejected wavelengths where $(D - \bar{D})$ were either lower than $-0.5\sigma_D$ or above $3\sigma_D$. This scheme was iterated ten times, which guarantees that we properly extracted the continuum envelope of the initial spectra. Finally, we use the remaining flux values to renormalize the initial spectra. As Fig. 3 shows, this concerns relatively small corrections of the continuum level that never exceed a few percent. As mentioned earlier by Muñoz Bermejo et al. (2013), this procedure may be debatable, but we found it satisfactory, and it was consistently applied to all the spectra we used.

3.2. Reduction of dimensionality

The ESPaDOnS and Narval spectropolarimeters typically provide about 250 000 flux measurements vs. wavelength across a spectral range spanning from about 380 to 1000 nm for each spectrum. We only consider spectra obtained in the polarimetric mode at a resolvance of $\mathcal{R} \sim 65\,000$. One of our main objective is to use stellar parameters derived from Stokes I spectra directly

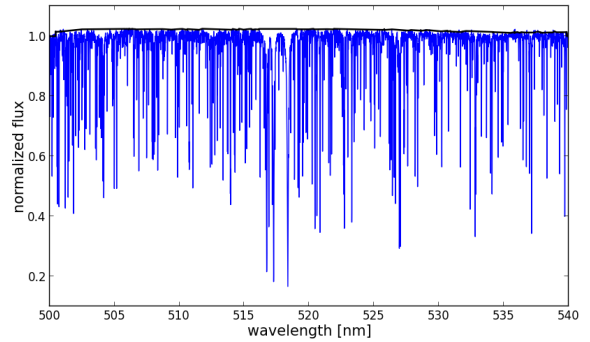


Fig. 3. Typical example of the continuum level fit (strong dark line) on top of the original spectrum that still has small normalization errors that need to be corrected for.

for the subsequent post-processing of the multi-line polarized spectra that are obtained simultaneously (see e.g., Paletou 2012 and references therein).

We adopted a first reduction of dimensionality by restricting the spectral domain from which we inverted stellar parameters to the vicinity of the Mg b-triplet, that is for wavelengths ranging from 500 to 540 nm. Good surface gravity indicators such as the strong lines of the b-triplet of Mg ($\lambda\lambda$ 516.75, 517.25, and 518.36 nm) and those of several other metallic lines in their neighbourhood, and the absence of telluric lines in this spectral domain are the main arguments we used. The matrix S that represents our training database size is $N_{\text{spectra}} = 905$ by $N_{\lambda} = 8\,000$.

Next, we computed the eigenvectors $e_k(\lambda)$ of the variance-covariance matrix defined as

$$C = (S - \bar{S})^T \cdot (S - \bar{S}), \quad (1)$$

where $\bar{S}$ is the mean of S along the N_{spectra} -axis. Therefore C is an $N_{\lambda} \times N_{\lambda}$ matrix. In the framework of the PCA reduction of dimensionality is reduced by representing the original data with a limited set of projection coefficients

$$p_{jk} = (S_j - \bar{S}) \cdot e_k, \quad (2)$$

with $k_{\text{max}} \ll N_{\lambda}$. For the processing of all observed spectra, we adopted $k_{\text{max}} = 12$.

⁵ We had to correct the value of 25 km s^{-1} given for HD 16232 because it is significantly underestimated compared with alternative determinations (with a median of 40 km s^{-1}).

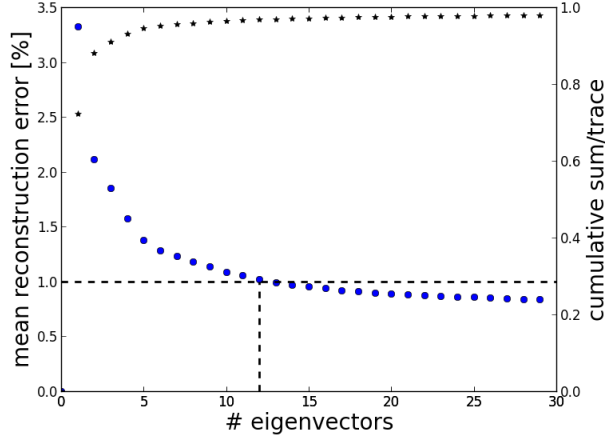


Fig. 4. Reconstruction error as a function of the number of eigenvectors used for computing of the projection coefficients p . The mean reconstruction error drops below 1% (with a standard deviation of 0.4%) after rank 12. The right-hand scale shows the cumulative sum of the ordered eigenvalues normalized to their sum ($\star$).

The most frequent argument supporting the choice of $k_{\max}$ relies on the accuracy achieved for reconstructing the original set of S_i 's from a limited set of eigenvectors (see e.g., Rees et al. 2000 or Muñoz Bermejo et al. 2013). In the present case, we display in Fig. 4 the mean reconstruction error

$$E(k_{\max}) = \left(\left| \frac{\bar{S} + \sum_{k=1}^{k_{\max}} p_{jk} \mathbf{e}_k - S_j}{S_j} \right| \right), \quad (3)$$

as a function of the maximum number of eigenvectors considered for computing the projection coefficients p . Figure 4 shows that this reconstruction error is better than 1% for $k_{\max} \geq 12$. In addition, at the same rank, an alternative estimator such as the cumulative sum of the ordered eigenvalues normalized to their sum becomes higher than 0.95.

To conclude this part, we display in Fig. 5 the 12 eigenvectors that we used below.

3.3. Nearest-neighbour(s) search

The above described reduction of dimensionality allows performing a fast and reliable inversion of observed spectra, after they have been (i) corrected for the wavelength shift vs. the spectra in our database, because of the radial velocity of the target; (ii) continuum-renormalized as accurately as possible; (iii) degraded in spectral resolution to be similar to the $\mathcal{R} \sim 42\,000$ resolvance of the Elodie spectra we use and; finally (iv) resampled in wavelength like the collection of Elodie spectra. We return to these various stages in the next section. After these tasks have been achieved, the inversion process is the following:

$O(\lambda)$ denotes the observed spectrum made similar to those of Elodie. We now compute the set of projection coefficients

$$\varrho_k = (O - \bar{S}) \cdot \mathbf{e}_k. \quad (4)$$

The nearest-neighbour search is made by seeking the minimum of the squared Euclidian distance

$$d_j^{(O)} = \sum_{k=1}^{k_{\max}} (\varrho_k - p_{jk})^2, \quad (5)$$

Table 1. Bias and standard deviation of the differences between inverted Elodie spectra and their initial stellar parameters of reference T_{eff} (K), $\log g$, $[\text{Fe}/\text{H}]$ (dex), and $v \sin i$ (km s^{-1}).

Parameter	T_{eff}	$\log g$	$[\text{Fe}/\text{H}]$	$v \sin i$
bias	0.78	0.02	0.005	0.07
σ	170	0.16	0.11	6.84

Notes. At each step, the Elodie spectra that was processed was removed from the training database.

where j spans the number, or a limited number if any a priori is known about the target, of distinct reference spectra in the training database. In practice, we did not limit ourselves to the nearest-neighbour search, although it already provides a relevant set of stellar parameters. Because PCA-distances between several neighbours may be of the same order, we adopted a simple procedure that consists of considering all neighbours in a domain

$$\min(d_j^{(O)}) \leq d_j^{(O)} \leq 1.2 \times \min(d_j^{(O)}), \quad (6)$$

and derived stellar parameters as the (simple) mean of each set of parameters $\{T_{\text{eff}}; \log g; [\text{Fe}/\text{H}]; v \sin i\}$ that characterises this set of nearest neighbours (A. López Ariste, priv. comm.). No significant changes in the results were recorded either for a smaller range of PCA-distances or when adopting, for example, distance-weighted mean parameters. This point is discussed again below.

3.4. Internal error

To characterise our inversion method, we inverted the stellar parameters T_{eff} , $\log g$, $[\text{Fe}/\text{H}]$ and $v \sin i$, for every spectrum (905) that constitutes the training database. However, at each step we removed the spectra that were processed from the database (and then recomputed the eigenvalues and eigenvectors of the new variance-covariance matrix).

To summarize this analysis, we list in Table 1 the internal errors, σ , measured for each inverted stellar parameter. The disappointing result on $v \sin i$ mainly comes from a suspicious scatter of results for these objects with large $v \sin i$, typically beyond 100 km s^{-1} . However, for our tests with S⁴N and P B data, we did not have to work with objects with $v \sin i$ beyond 80 km s^{-1} (see next sections).

4. Inversion of S⁴N spectra

A convincing test of our method would be to invert high-resolution spectra that were gathered in the frame of the S⁴N spectroscopic survey (Allende Prieto et al. 2004).

We only selected the S⁴N spectra for which all three parameters $\{T_{\text{eff}}; \log g; [\text{Fe}/\text{H}]\}$ have been determined by Allende Prieto et al. (2004). We used the same reference catalogue as the one used for the Elodie spectra to add a $v \sin i$ value to each spectrum (see e.g., Paletou & Zolotukhin 2014). Since they were acquired at a higher spectral resolution than those of Elodie, we adapted each spectrum to the resolvance of Elodie of $\mathcal{R} \sim 42\,000$ using an appropriate Gaussian filter. Finally we applied the same renormalization procedure to all spectra as was applied to our Elodie spectra database.

We identified 49 objects in common between S⁴N and our sample of 905 objects taken from the Elodie stellar library. Using the catalogue values for effective temperature T_{eff} , surface gravity $\log g$, and metallicity $[\text{Fe}/\text{H}]$, we easily estimated the bias and

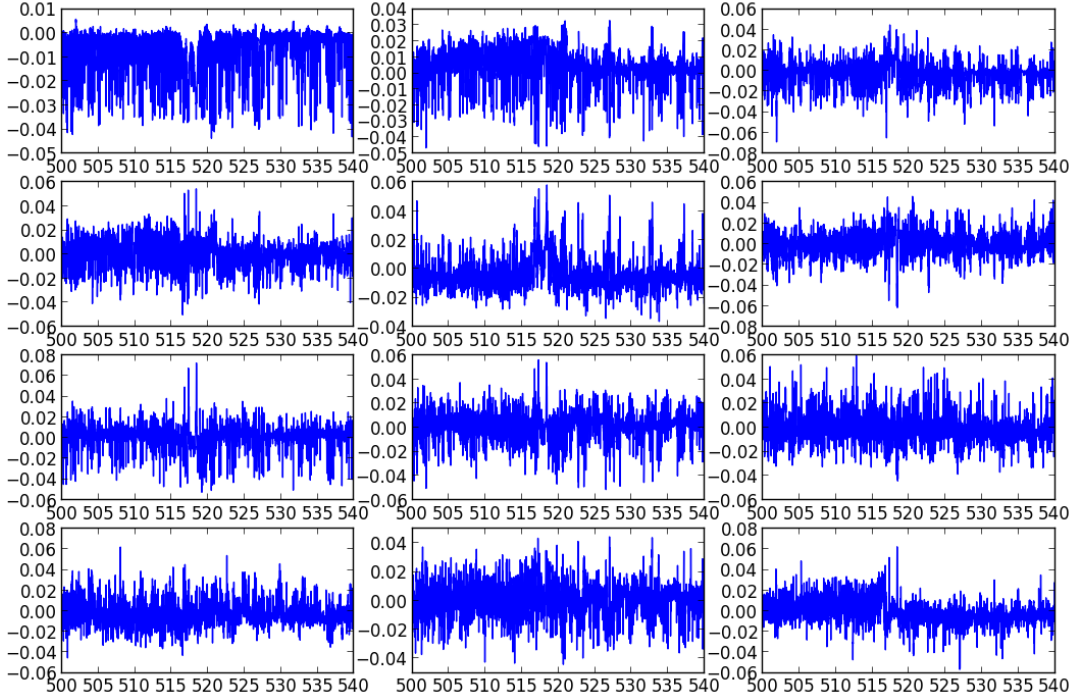


Fig. 5. From top left to bottom right, the 12 eigenvectors used in our inversion method. Each of them is displayed vs. the same wavelength scale in nm.

Table 2. Bias and standard deviation of the absolute differences between reference stellar parameters T_{eff} (K), $\log g$ (dex), and $[\text{Fe}/\text{H}]$ (dex) retrieved from the S⁴N and the Elodie catalogues for the 49 objects in common between our data samples.

Parameter	T_{eff}	$\log g$	$[\text{Fe}/\text{H}]$
bias	-66	0.16	0.006
σ	84	0.14	0.08

Table 3. Bias and standard deviation of the absolute differences between inverted and reference stellar parameters T_{eff} (K), $\log g$, $[\text{Fe}/\text{H}]$ (dex), and $v \sin i$ (km s⁻¹) for our sample of 104 S⁴N spectra – see also Fig. 6.

Parameter	T_{eff}	$\log g$	$[\text{Fe}/\text{H}]$	$v \sin i$
bias	-74	0.16	-0.01	-0.36
σ	115	0.16	0.11	1.61
49 objects in common				
bias	-62	0.16	0.001	-0.16
σ	82	0.16	0.08	0.67
55 objects not in common				
bias	-85	0.15	-0.02	-0.55
σ	138	0.17	0.13	2.10

standard deviations between the two distinct estimate for each stellar parameter. The results are summarized in Table 2.

These results were then compared with the results of our inversion of 104 S⁴N spectra using our Elodie training database for PCA. They are summarized, for each stellar parameter including $v \sin i$ in Fig. 6 and Table 3 (where we also detail specific values for the inverted spectra of the 49 objects in common between S⁴N and our Elodie sample and the 55 remaining objects). Bias and dispersions, especially conspicuous for $\log g$, measured after our inverted parameters appear to be quite direct imprints

of the discrepancies that were already clear from the direct comparison between the reference values for objects in common between the two samples. However, taking into account that our inversion only relies on spectroscopic information from a limited (but relevant) spectral band, we find our approach satisfactory for our purpose.

Our method also allows a reliable and direct inversion of the projected rotational velocity $v \sin i$ of stars, without any other limitation than the one coming from the limit in $v \sin i$ attached to the identified values for spectra present in our Elodie training database – see Figs. 2, for example. Using synthetic spectra (computed for hypothetical non-rotating stars), Gazzano et al. (2010) were limited to stars with a $v \sin i$ lower than about 11 km s⁻¹, for instance. In our case, we accurately retrieved $v \sin i$ values up to the most extreme case of ~ 80 km s⁻¹ in our sample that is, the high-proper motion F1V star HIP 77952⁶ that is, fast rotators for which usual synthetic models are unsatisfactory.

Another advantage of our method is that nearest-neighbour(s) are also identifiable as objects that is, as other stars. In that sense, our method can also be seen as relevant for classification. For instance, considering the only objects in common between our S⁴N and our Elodie spectra and objects sample, in 75% of the cases the very nearest neighbour is another spectrum of the same object and, for the remainder, conditions expressed in Eq. (6) guarantee that the same object spectrum belongs to the set of nearest neighbours.

5. Inversion of PolarBase spectra

Below we first discuss details of the kind of conditioning that has to be applied to P B spectra to make them similar to

⁶ Note that for this object, five estimates of $v \sin i$ from VizieR at CDS can be identified that range from 70 to 90 km s⁻¹, but whose median value of 75 km s⁻¹ agrees very well with our own estimate.

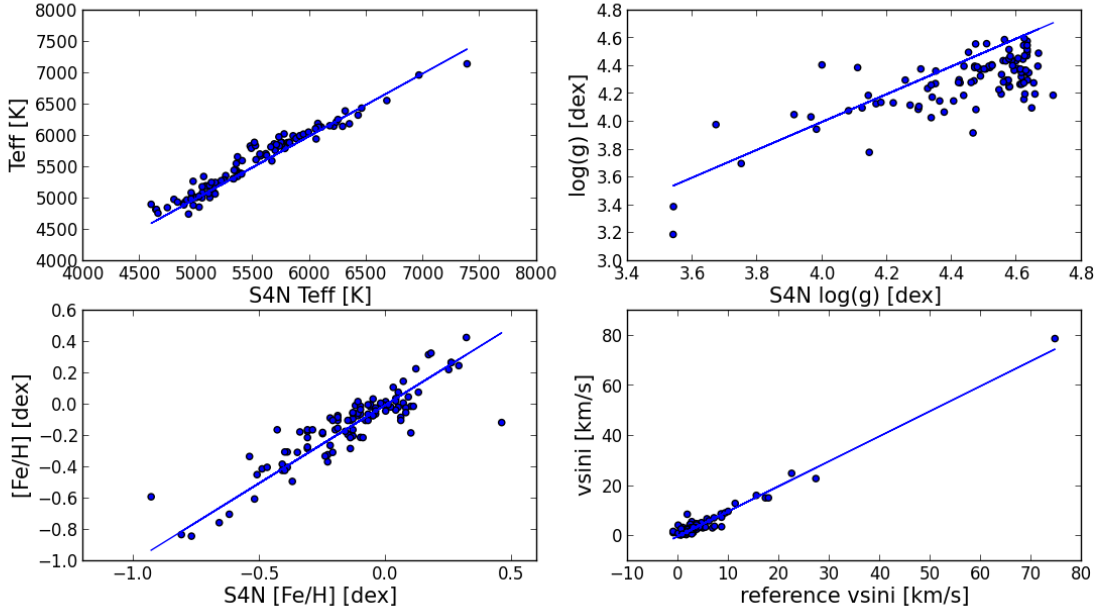


Fig. 6. Stellar fundamental parameters inverted from the analysis of S⁴N spectra, using our PCA-Elodie method vs. their S⁴N catalogue values of reference (except for $v \sin i$ – see text). The bias and standard deviation deduced from the analysis of the difference between inverted and catalogue and reference parameters are $(-74; 115)$ for T_{eff} , $(0.16; 0.16)$ for $\log g$ and $(-0.01; 0.11)$ for $[\text{Fe}/\text{H}]$. For $v \sin i$ values for which the source of data is not the S⁴N catalogue, however, we find $(-0.36; 1.61)$.

our Elodie-based training database. Then we discuss the inversion of solar spectra taken on reflection over the Moon surface at the TBL telescope, and finally inversions of 140 spectra from FGK objects in common with the sample studied by Valenti & Fischer (2005).

5.1. Conditioning of PolarBase spectra

The first and obvious task to perform on observed spectra is to correct for their wavelength shift vs. Elodie spectra, which are found already corrected for radial velocity (v_{rad}). The radial velocity of the target at the time of the observation is deduced from the centroid, in a velocity space, of the pseudo-profile resulting from the addition (see e.g., Paletou 2012) of the three spectral lines of the Ca infrared triplet whose rest wavelengths are 849.802, 854.209 and 866.214 nm. This can be done with any other set of spectral lines assumed to be a priori present in the spectra we wish to process. One of the advantages of the Ca infrared triplet is its persistence for spectral types ranging from A to M. When v_{rad} is known, the observed profile is set on a new wavelength grid, at rest.

We checked with a solar spectrum (see also the next section) that v_{rad} have to be known to an accuracy of the order of $\delta v/4$ with $\delta v = c/\mathcal{R}$ i.e., about 1.15 km s^{-1} with our ESPaDOnS-Narval data. Beyond this value, estimates of T_{eff} and $v \sin i$ first start to be significantly affected by the misalignment of the observed spectral lines with those of the spectra of the training database. Indeed, the neighbourhood identified by our PCA-based approach can change quite dramatically because of such a spectral misalignment.

A second step consists of adapting the resolution of the initial spectra, about $\mathcal{R} \sim 65\,000$ in the polarimetric mode, to the resolution of the training database spectra that is, $\mathcal{R} \sim 42\,000$. This is done by convolving the initial observed profile by a Gaussian profile of adequate width. Then we resampled the wavelength

grid down to the one common to all reference spectra and we interpolated the original spectra onto the new wavelength grid. Finally, we applied the same renormalization procedure as described in Sect. 3.1, for consistency.

5.2. Solar spectra observed with Narval

First tests of our inversion method with P B data were performed using solar spectra observed by the 2 m aperture TBL telescope by reflection over the surface of the Moon in March and June 2012.

Using the very same training database as was used for the tests done with S⁴N data, we identified as nearest-neighbour star to the Sun HD 186427 (also known as 16 Cyg B). It is a G3V planet-hosting star often identified as a solar twin (see e.g., Porto de Mello et al. 2014). Its stellar fundamental parameters are $T_{\text{eff}} = 5757 \text{ K}$, $\log g = 4.35 \text{ dex}$, $[\text{Fe}/\text{H}] = 0.06$, and $v \sin i \sim 2.18 \text{ km s}^{-1}$ (see also Tucci Maia et al. 2014, for a recent determination of these parameters). The other nearest neighbour we identified is HD 29150, a star whose main parameters are $T_{\text{eff}} = 5733 \text{ K}$, $\log g \sim 4.35 \text{ dex}$ (Lee et al. 2011, from Simbad at CDS query), $[\text{Fe}/\text{H}] = 0.0$, and $v \sin i \sim 1.8 \text{ km s}^{-1}$.

Taking into consideration this neighbourhood, in the PCA-sense, we thus derived quite satisfactory estimates for the effective temperature $T_{\text{eff}} = 5745 \text{ K}$, surface gravity $\log g = 4.35 \text{ dex}$, a metallicity of $[\text{Fe}/\text{H}] = 0.03$, and $v \sin i \sim 2 \text{ km s}^{-1}$ typical of the (very) slowly rotating Sun, as observed by reflection over the Moon surface with the Narval at TBL spectropolarimeter.

This is quite consistent with the test consisting of inverting solar spectra taken from the Elodie stellar library, but not a member of our training database. In this case, we recover neighbours 16 Cyg B and HD 29150 plus the additional HD 146233 (also known as 18 Sco) and HD 42807, a RSCVn star of G2V type, also very similar to the Sun.

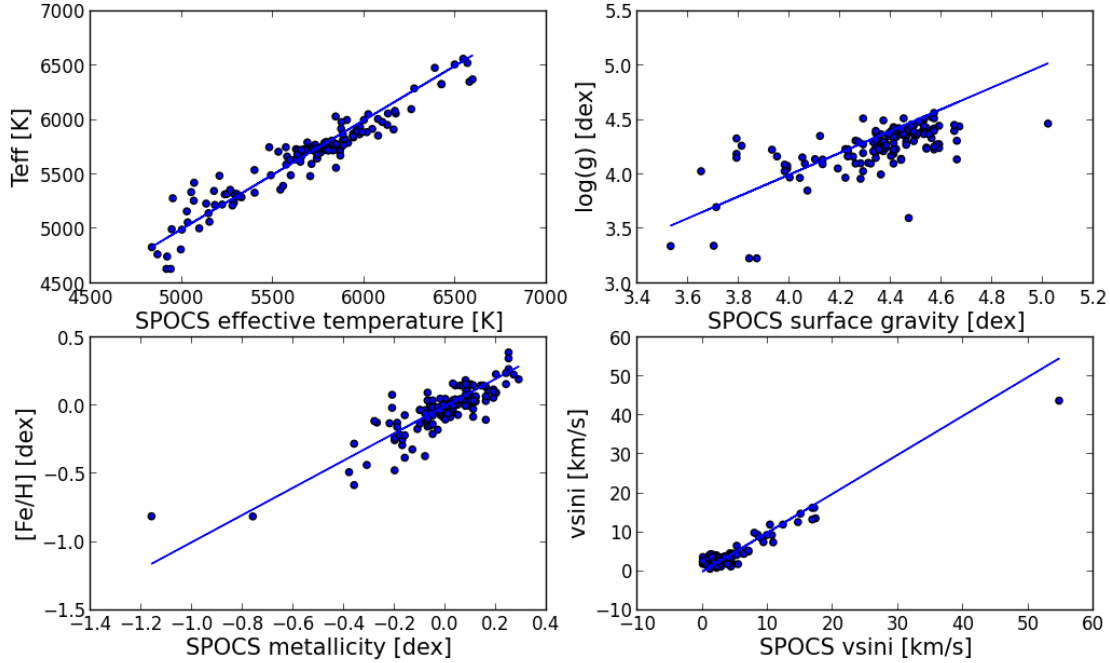


Fig. 7. Stellar fundamental parameters inverted from P B spectra in common with the S catalogue, using our PCA-Elodie method vs. S reference values, including $v \sin i$. Respective bias and standard deviation deduced from the analysis of the difference between inverted and reference parameters are (23; 115) for T_{eff} , (0.10; 0.19) for $\log g$, (0.02; 0.10) for $[\text{Fe}/\text{H}]$, and $(-0.04; 1.68)$ for $v \sin i$.

Table 4. Bias and standard deviation of the absolute differences between stellar parameters T_{eff} (K), surface gravity $\log g$, metallicity $[\text{Fe}/\text{H}]$ (dex), and projected rotational velocity $v \sin i$ (km s^{-1}) obtained from the S catalogue of Valenti & Fischer (2005, reference values) and our inversion method using Elodie spectra.

Parameter	T_{eff}	$\log g$	$[\text{Fe}/\text{H}]$	$v \sin i$
bias	23	0.10	0.02	-0.04
σ	115	0.19	0.10	1.68

5.3. Other FGK stars

We identified in the present content of P B 140 targets that are also identified in the S catalogue. For our next tests of our method, we selected the spectra with the best signal-to-noise ratio for every object. Typical values were already given in Fig. 1.

Results and characterization of our inversion are given in Table 4 and with more details in Fig. 7. This figure also gives an idea about the range of variations of parameters that are expected for the set of spectra and objects that are studied here. Overall, the figures are quite similar to those obtained with the S⁴N, although bias values for T_{eff} and $\log g$ are lower for the S data.

5.3.1. Effective temperature

This standard deviation on the differences between our values and those of the S reference is the same as the difference we evaluated from S⁴N spectra (see Table 3). However, the bias value of 23 K is much lower for this sample of objects and P B spectra. We note also that the most important dispersion is for the coolest objects of our sample with a T_{eff} of about 5000 K.

A more detailed inspection of outliers in effective temperature, taking into account alternative and more recent estimates of T_{eff} than the estimate adopted from S, reveals that the most extreme ΔT_{eff} we identified are quite often overestimated. This is for instance the case for the K1V star LHS 44, for which a recent determination by Maldonado et al. (2012) is +200 K higher than that of Valenti & Fischer (2005) and only 140 K higher than ours. Another effect may also come from new estimates of parameters for Elodie objects themselves, an issue we shall discuss below.

Our estimates of the effective temperature from spectropolarimetric data are satisfactory and can be used in selecting a proper mask (i.e., at least a list of spectral lines) that will be, in turn, used for the subsequent extraction of polarized signatures.

5.3.2. Surface gravity

It is well known that surface gravity is the most difficult parameter to extract from the analysis of spectroscopic data. The most conspicuous outliers are quite easily detected on the $\log g$ subplot in Fig. 7. The top three of them, which show $\Delta \log g \sim 0.5$ dex or higher, we find EK Dra, LTT 8785 and HD 22918. For EK Dra, our estimate of $\log g \sim 3.6$ dex is much too low compared with the (rare) values found in the literature (~ 4.5 dex). This may be compensated by the fact that its identified nearest-neighbour star is the G2IV subgiant HD 126868 (also known as 105 Vir), for which $\log g = 3.6$ dex in the Elodie catalogue, although a value of 3.9 dex is reported elsewhere (see e.g., the P catalogue: Soubiran et al. 2010). For LTT 8785 and HD 22918, an examination of alternative estimates for $\log g$ (e.g., Massarotti et al. 2008; Jones et al. 2011) shows that S values may have been slightly overestimated. In addition, the surface gravity ($\log g \sim 3.23$) of the nearest neighbour (and the

same object in both cases) HD 42983 may have been underestimated. An inspection of VizieR at CDS for this object indicates four different estimates ranging from 3.23 to 3.6 with a median value of 3.5. Taking this into account, $\Delta \log g$ does not exceed 0.2 dex between our inverted values and reference ones, which is satisfactory.

5.3.3. Metallicity

We now inspect our results for metallicity. First of all, we included LHS 44 in our sample, which has $[\text{Fe}/\text{H}] = -1.16$ according to Valenti & Fischer (2005), a value a priori excluded from our working range. But our inversion method still indicates these objects in our sample as bearing the lowest $[\text{Fe}/\text{H}]$ values. Then by decreasing order of $\Delta[\text{Fe}/\text{H}]$ of our inverted values and the reference values, we find HD 30508, 40 Eri, and LHS 3976. We found a systematically better agreement between our estimate and statistics on all data available at VizieR, to within 0.05 dex.

5.3.4. Projected rotational velocity

Our determinations of $v \sin i$ are correct and especially interesting for objects with a significant projected rotational velocity, for example, greater than about 10 km s^{-1} . The main outlier, as seen in Fig. 7, is HR 1817, for which we derive a $v \sin i \sim 43 \text{ km s}^{-1}$ while the SPOCS value is about 55 km s^{-1} . It is a F8V RS CVn star that is also known as AF Lep, for which another value of 52.6 km s^{-1} , still about 10 km s^{-1} greater than our estimate, was more recently published by Schröder et al. (2009).

Other methods for determining $v \sin i$ already exist. However, to the best of our knowledge, they require a template (synthetic) spectrum at $v \sin i \sim 0$ or, at least, a list of spectral lines a priori expected in the spectra, as auxiliary and support data (see e.g., Díaz et al. 2011 and references therein). Data-processing tools that we attach to P-B will therefore include a complementary Fourier analysis module providing an additional $v \sin i$ determination, once stellar fundamental parameters will be available from our inversion tool⁷. Note also that with our PCA-based method, we are mostly interested in the intermediate $v \sin i$ regime, that is between 10 and 100 km s^{-1} . Indeed, for slower rotators for which rotational broadening becomes of the order of other sources of broadening (e.g., instrumental or turbulent), a more detailed or specific line profile analysis may be required.

6. Discussion

As briefly remarked earlier, the fact that the current implementation of our method relies on observed data makes it also somewhat relevant to classification. Indeed, we did not only identify nearest spectra since nearest neighbour(s) can also be identified as objects that is, as other stars. This important fact is also totally independent of the method used to determine the stellar parameters for these objects.

Therefore, unless modifying the sample of spectra and objects that constitute our training database, we do not expect any change in the relation between inverted spectra and nearest neighbours as objects, even though evaluations of their various stellar parameters may still change in time. Another interesting point is that this is also true for any other stellar parameter – especially those contributing to the spectral signature of a star,

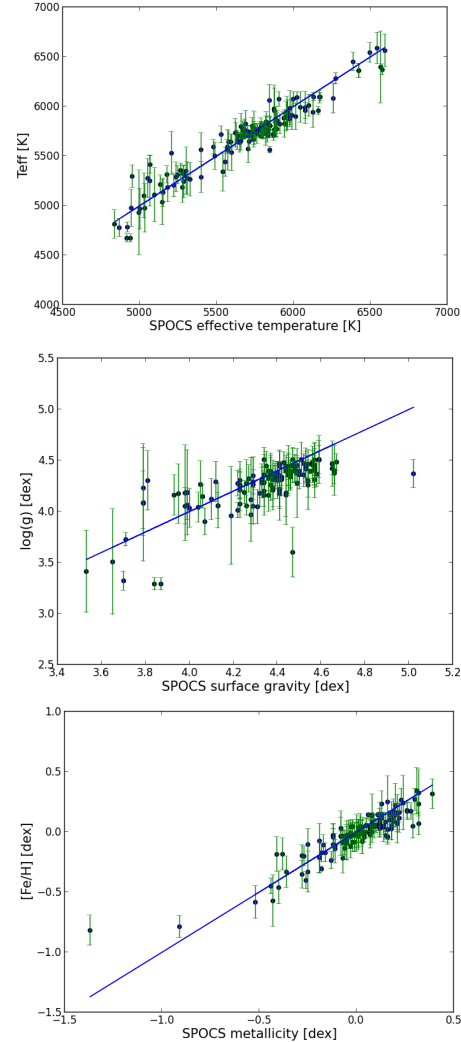


Fig. 8. Stellar fundamental parameters and their respective uncertainties inverted from P-B spectra in common with the SPOCS catalogue, obtained using our PCA-Elodie method vs. SPOCS reference values. These values were deduced from all determinations for each nearest neighbour compiled in the P-B catalogue of Soubiran et al. (2010).

beyond the limited set of fundamental parameters we considered in this study.

Taking advantage of this, instead of using values given by the Elodie catalogue alone, for each SPOCS spectra we analysed, we gathered for every nearest neighbour all the evaluations of effective temperature, surface gravity, and metallicity provided by the comprehensive P-B catalogue⁸ (Soubiran et al. 2010). We also evaluated uncertainties on each fundamental parameter as the standard deviation of the full set of collected values, except for the unavailable projected rotational velocity. The results are displayed in Figs. 8, and details are given in Table 5. Mean errors given in Table 5 are 110 K, 0.16 dex and 0.09 dex for T_{eff} , $\log g$ and $[\text{Fe}/\text{H}]$ respectively which is quite consistent, but slightly better than the standard deviation values already given in Table 4.

⁷ The same is true for the refinement of radial velocity measurements, which can be improved by line addition when a proper list of expected spectral line wavelengths (at rest) is identified.

⁸ For four targets for which we could not retrieve data from P-B, we instead used instead their TGMET values (Katz et al. 1998) that are also available from the Elodie at VizieR catalogue.

Table 5. Values and uncertainties of the inverted stellar parameters T_{eff} , $\log g$, and $[\text{Fe}/\text{H}]$ for S objects in common with P B using our Elodie-based PCA method.

Object	T_{eff} [K]	$\sigma_{T_{\text{eff}}}$ [K]	$\log g$ [dex]	$\sigma_{\log g}$ [dex]	$[\text{Fe}/\text{H}]$ [dex]	$\sigma_{[\text{Fe}/\text{H}]}$ [dex]
HR753	4777	91	4.40	0.24	-0.03	0.09
LTT11292	5562	203	4.46	0.14	-0.02	0.11
*39Tau	5763	75	4.36	0.12	0.05	0.06
LTT11169	5098	226	3.51	0.52	-0.10	0.13
LTT11282	5990	153	4.21	0.20	0.04	0.10
*109Psc	5768	100	4.29	0.20	0.17	0.10
*107Psc	5183	145	4.45	0.17	-0.03	0.12
LHS5051a	5645	86	4.40	0.03	0.13	0.03
HD12846	5734	100	4.31	0.16	-0.20	0.09
LTT10989	5702	19	4.39	0.02	0.14	0.01
*13Tri	5854	77	3.90	0.13	-0.24	0.07
HD30825	5312	85	3.73	0.07	-0.11	0.08
LHS1753	5821	98	4.26	0.15	0.03	0.07
G175-33	5724	120	4.21	0.17	0.23	0.09
HD9472	5718	39	4.48	0.08	-0.00	0.05
V*V451And	5732	144	4.47	0.10	-0.00	0.08
LTT10580	5702	19	4.39	0.02	0.14	0.01
LHS1125	4974	242	4.51	0.14	0.05	0.12
HD5065	5945	154	4.05	0.22	-0.10	0.14
*etaCas	5825	140	4.18	0.23	-0.33	0.17
LTT11104	6369	53	3.97	0.09	0.08	0.06
V*V987Cas	5264	168	4.46	0.14	0.00	0.12
G245-27	5649	91	4.42	0.14	0.04	0.09
HD73344	6095	41	4.18	0.13	0.11	0.04
*55Cnc	5287	151	4.43	0.09	0.34	0.19
LTT12401	5897	134	4.30	0.16	-0.00	0.09
HR3499	6080	150	4.18	0.27	0.04	0.12
*24LMi	5768	83	4.26	0.13	0.05	0.05
LHS2216	5749	50	4.35	0.18	0.16	0.06
HR4486	5840	154	4.36	0.17	-0.21	0.13
*36UMa	5940	156	4.20	0.24	-0.22	0.12
HD98618	5773	57	4.34	0.11	0.04	0.05
V*V377Gem	5781	168	4.48	0.07	-0.02	0.08
NLTT17627	5804	86	4.34	0.12	-0.01	0.09
LTT12204	5822	81	4.34	0.13	-0.02	0.07
HD71881	5845	96	4.26	0.15	-0.03	0.08
HD138573	5739	65	4.36	0.11	0.04	0.06
HR5659	5702	19	4.39	0.02	0.14	0.01
GJ566A	5591	69	4.42	0.07	-0.03	0.04
HD129814	5808	88	4.31	0.13	-0.01	0.08
*tauBoo	6449	93	4.27	0.23	0.26	0.10
*sigBoo	6396	362	3.96	0.48	-0.57	0.21
LHS348	5987	163	4.31	0.18	0.05	0.09
LHS2498	5275	195	4.09	0.57	-0.19	0.13
V*EKDra	5561	32	3.60	0.24	-0.04	0.02
V*HNPEG	5872	78	4.51	0.12	-0.09	0.08
HR8455	5796	129	4.17	0.29	-0.11	0.13
V*V376Peg	6009	112	4.23	0.21	-0.04	0.11
HD195019	5787	96	4.25	0.13	0.05	0.07
LTT16813	5341	198	4.40	0.25	0.07	0.16
V*V454And	5710	54	4.44	0.13	0.08	0.08
HR8832	4812	142	4.46	0.18	0.10	0.15
LHS544	5213	92	4.39	0.19	0.00	0.10
LTT15779	5770	69	4.32	0.12	0.06	0.05
LTT15881	5702	19	4.39	0.02	0.14	0.01
LTT16028	5805	33	4.39	0.15	-0.11	0.05
HD173701	5286	156	4.43	0.10	0.32	0.20
HR7294	5745	70	4.36	0.10	0.04	0.06
LTT15729	6562	164	4.04	0.14	-0.04	0.16
*16CygB	5776	70	4.28	0.11	0.06	0.05
*sigDra	5310	89	4.45	0.14	-0.07	0.14

Table 5. continued.

Object	T_{eff} [K]	$\sigma_{T_{\text{eff}}}$ [K]	$\log g$ [dex]	$\sigma_{\log g}$ [dex]	[Fe/H] [dex]	$\sigma_{[\text{Fe}/\text{H}]}$ [dex]
LTT149	5754	62	4.33	0.16	0.16	0.09
LHS146	5319	81	4.51	0.24	-0.58	0.14
HD12328	4782	50	3.32	0.09	0.05	0.10
HR448	5881	57	4.06	0.12	0.18	0.07
LTT1267	5778	64	4.27	0.12	0.10	0.07
HD3821	5738	98	4.36	0.13	-0.14	0.10
G270-82	5634	98	4.40	0.15	-0.21	0.13
LTT709	6090	57	4.19	0.11	0.05	0.06
LTT10409	5412	90	4.30	0.29	-0.19	0.14
HR222	4976	186	4.51	0.17	-0.20	0.08
HD9986	5744	66	4.36	0.11	0.04	0.06
LTT8887	5440	148	4.28	0.34	0.15	0.15
LTT9062	5764	87	4.32	0.15	-0.03	0.10
HR8931	5884	47	4.05	0.17	-0.40	0.07
G157-8	5897	147	4.29	0.24	-0.14	0.12
HR8734	5570	137	4.26	0.24	0.32	0.12
LTT8785	4672	38	3.29	0.06	-0.05	0.05
HD208776	5981	172	4.03	0.19	-0.07	0.16
HD377	5975	236	4.12	0.41	-0.03	0.03
LHS1239	5724	99	4.35	0.14	-0.17	0.08
G131-59	5799	86	4.32	0.13	-0.01	0.09
*54Psc	5204	119	4.36	0.19	0.02	0.12
HD218687	6073	72	4.45	0.14	-0.04	0.04
LTT16778	5348	238	4.18	0.42	0.22	0.18
V*MTPEG	5783	88	4.36	0.14	0.02	0.07
*51Peg	5702	19	4.39	0.02	0.14	0.01
G131-18	5355	72	4.51	0.18	-0.45	0.07
V*V439And	5640	110	4.37	0.12	0.03	0.09
V*V344And	5660	96	4.37	0.10	0.07	0.09
LHS3976	5666	79	4.31	0.18	-0.46	0.14
V*AFLeP	6542	99	4.37	0.14	0.05	0.10
HR2622	5957	81	4.05	0.09	0.09	0.05
HD46375	5243	164	4.42	0.13	0.24	0.23
*alfCMi	6587	157	4.06	0.14	-0.01	0.15
LTT2093	5882	60	4.07	0.13	0.17	0.07
*40Eri	5133	142	4.45	0.22	-0.35	0.11
HD30508	5528	215	4.08	0.32	0.03	0.14
HR1232	4928	426	3.41	0.40	0.11	0.12
LTT1723	5249	258	4.23	0.40	0.13	0.18
LHS1577	5823	134	4.22	0.19	-0.79	0.09
HD22918	4672	38	3.29	0.06	-0.05	0.05
V*kap01Cet	5745	101	4.45	0.09	0.08	0.11
*1Ori	6361	73	4.05	0.18	-0.08	0.10
HR2251	5857	99	4.28	0.15	-0.02	0.08
LTT11933	5911	151	4.35	0.20	-0.09	0.10
*37Gem	5770	82	4.16	0.16	-0.31	0.02
V*chi01Ori	5872	78	4.51	0.12	-0.09	0.08
LTT5873	5275	164	4.46	0.14	-0.01	0.11
LTT3686	5702	19	4.39	0.02	0.14	0.01
LHS2465	5956	40	4.01	0.07	0.10	0.03
LTT12723	5751	67	4.38	0.10	0.04	0.05
*17Vir	6095	41	4.18	0.13	0.11	0.04
V*LWCom	5702	19	4.39	0.02	0.14	0.01
LTT13145	5910	146	4.31	0.19	-0.03	0.11
LTT13442	6281	57	4.15	0.11	0.17	0.07
*61UMa	5500	137	4.48	0.10	-0.04	0.10
LHS44	5294	112	4.47	0.22	-0.82	0.12
LTT7713	5762	61	4.32	0.07	0.06	0.05
HD175726	6073	72	4.45	0.14	-0.04	0.04
*18Sco	5752	75	4.35	0.13	0.04	0.06
V*V2133Oph	5183	150	4.44	0.18	-0.02	0.12
V*V2292Oph	5591	69	4.42	0.07	-0.03	0.04
HD169822	5643	73	4.44	0.10	-0.04	0.07

Table 5. continued.

Object	T_{eff} [K]	$\sigma_{T_{\text{eff}}}$ [K]	$\log g$ [dex]	$\sigma_{\log g}$ [dex]	[Fe/H] [dex]	$\sigma_{[\text{Fe}/\text{H}]}$ [dex]
HD159909	5733	73	4.35	0.13	0.08	0.08
HR6950	5562	169	4.16	0.20	0.23	0.11
*110Her	6361	73	4.05	0.18	-0.08	0.10
LTT15317	5747	77	4.40	0.13	-0.09	0.10
HD166435	6060	160	4.16	0.29	-0.02	0.07
*86Her	5535	184	4.12	0.20	0.27	0.09
HR6806	4966	205	4.45	0.17	-0.17	0.08
V*delEri	5108	270	4.18	0.47	0.25	0.21
V*epsEri	5034	228	4.51	0.16	-0.12	0.10
LTT1601	5716	84	4.26	0.23	-0.33	0.13
LHS1845	5748	70	4.35	0.10	0.04	0.06
V*V401Hya	5746	63	4.39	0.11	0.08	0.08
LTT3283	5702	19	4.39	0.02	0.14	0.01
HD145825	5742	71	4.37	0.12	0.04	0.06
HD143006	5956	227	4.30	0.18	0.16	0.08
HR7291	6094	60	4.16	0.11	0.05	0.07

As a final remark, it is also worth mentioning that although we used it together with a training database made of observed spectra, our PCA-based inversion method can be equally implemented using synthetic spectra.

7. Conclusion

We have implemented a fast and reliable PCA-based numerical method for inverting stellar fundamental parameters T_{eff} , $\log g$, and [Fe/H] and, the projected rotational velocity $v \sin i$, from the analysis of high-resolution Echelle spectra delivered by the Narval and ESPaDOnS spectropolarimeters. First tests, mainly made with FGK-star spectra, show a good agreement between our inverted stellar parameters and reference values published by Allende Prieto et al. (2004) and Valenti & Fischer (2005). We also believe that our method will also be efficient for analysing spectra from cooler M stars as well as from hotter stars, up to spectral type A.

We used it, so far, with a spectral band located at the vicinity of the b-triplet of Mg and without any help from additional (e.g., photometric) information, which is particularly challenging. However, we can easily either extend or change the spectral domain of use, or combine analyses from several distinct spectral domains to further constrain and refine the stellar parameters determination. In that respect, it is important to realize that PCA allows for a quite dramatic reduction of dimensionality, of the order of 800 ($\sim N_{\lambda}/k_{\text{max}}$) for the configuration we presented here. This capability is indeed of great interest for a comprehensive post-processing of high-resolution spectra covering a very broad bandwidth such as those from the ESPaDOnS and Narval spectropolarimeters.

Acknowledgements. This research has made use of the VizieR catalogue access tool, CDS, Strasbourg, France. The original description of the VizieR service was published in A&AS 143, 23. This research has made use of the SIMBAD database, operated at CDS, Strasbourg, France. P B data were provided by the OV-GSO (ov-gso.irap.omp.eu) datacenter operated by CNRS/INSU and the Université Paul Sabatier, Observatoire Midi-Pyrénées, Toulouse (France).

References

- Allende Prieto, C., Barklem, P. S., Lambert, D. L., & Cunha, K. 2004, A&A, 420, 183
- Bailer-Jones, C. A. L. 2000, ApJ, 357, 197
- Bailer-Jones, C. A. L., Irwin, M., & von Hippel, T. 1998, MNRAS, 298, 361
- Cabanac, R. A., de Lapparent, V., & Hickson, P. 2002, A&A, 389, 1090
- Casini, R., Asensio Ramos, A., Lites, B. W., & López Ariste, A. 2013, ApJ, 773, 180
- Cayrel, G., & Cayrel, R. 1963, ApJ, 137, 431
- Cayrel de Strobel, G. 1969, in Proc. of the 3rd Harvard-Smithsonian Conf. on Stellar Atmospheres, ed. O. Gingerich, 35
- Deeming, T. J. 1964, MNRAS, 127, 493
- Díaz, C. G., González, J. F., Levato, H., & Grosso, M. 2011, A&A, 531, A143
- Gazzano, J.-C., de Laverny, P., Deleuil, M., et al. 2010, A&A, 523, A91
- Głebocicki, R., & Gnaniński, P. 2005, VizieR On-line Data Catalog: III/244
- Jones, M. I., Jenkins, J. S., Rojo, P., & Melo, C. H. F. 2011, A&A, 536, A71
- Katz, D., Soubiran, C., Cayrel, R., Adda, M., & Cautain, R. 1998, A&A, 338, 151
- Lee, Y.S., Beers, T.C., Allende Prieto, C., et al. 2011, AJ, 141, 90
- Maldonado, J., Eiroa, C., Villaver, E., Montesinos, B., & Mora, A. 2012, A&A, 541, A40
- Massarotti, A., Latham, D. W., Stefanik, R. P., & Fogel, J. 2008, AJ, 135, 209
- Molenda-Zakowicz, J., Sousa, S. G., Frasca, A., et al. 2013, MNRAS, 434, 1422
- Munari, U. 1999, Baltic Astron., 8, 73
- Munari, U., Sordo, R., Castelli, F., & Zwitter, T. 2005, A&A, 442, 1127
- Muñoz Bermejo, J., Asensio Ramos, A., & Allende Prieto, C. 2013, A&A, 553, A95
- Paletou, F. 2012, A&A, 544, A4
- Paletou, F., & Zolotukhin, I. 2014 [[arXiv:1408.7026](https://arxiv.org/abs/1408.7026)]
- Petit, P., Louge, T., Théado, S., et al. 2014, PASP, 126, 469
- Porto de Mello, G. F., da Silva, R., da Silva, L., & de Nader, R. V. 2014, A&A, 563, A52
- Prugniel, P., & Soubiran, C. 2001, A&A, 369, 1048
- Prugniel, P., Soubiran, C., Koleva, M., & Le Borgne D. 2007 [[arXiv:astro-ph/0703658](https://arxiv.org/abs/astro-ph/0703658)]
- Recio-Blanco, A., Bijaoui, A., de Laverny, P., et al. 2006, MNRAS, 370, 141
- Rees, D. E., López Ariste, A., Thatcher, J., & Semel, M. 2000, A&A, 355, 759
- Re Fiorentin, P., Bailer-Jones, C. A. L., Lee, et al. 2007, A&A, 465, 1373
- Schröder, C., Reiners, A., & Schmitt, J. H. M. M. 2009, A&A, 493, 1099
- Semel, M., Donati, J. F., & Rees, D. E. 1993, A&A, 278, 231
- Soubiran, C., Le Campion, J.-F., Cayrel de Strobel, G., & Caillou, A. 2010, A&A, 515, A111
- Tucci Maia, M., Meléndez, J., & Ramírez, I. 2014, ApJ, 790, L25
- Valenti, J. A., & Fischer, D. A. 2005, ApJS, 159, 141

Chapitre 4

Apport de la régression inverse par tranches (SIR)

Dans le chapitre précédent, nous avons vu, grâce à l'application de l'ACP, comment nous pouvions utiliser une approche non-supervisée pour apporter une solution au problème d'estimation de paramètres stellaires fondamentaux. Les méthodes de réduction de dimension de type non-supervisées sont surtout intéressantes pour des aspects de compression ou de transmission (par exemple en télécommunications), dans lesquelles il faut minimiser l'entropie¹ de l'information. Dans le cas qui nous intéresse, ce sont les aspects de réduction du bruit (d'extraction de l'information pertinente) qui nous sont utiles (toujours au sens de "composante non-pertinente" du signal). Une approche supervisée pourrait maximiser cette mise en évidence de "l'information pertinente" (celle qui va apporter le plus de précision sur la valeur de paramètre que l'on recherche, en l'occurrence il s'agit des paramètres stellaires fondamentaux). SIR (Li, 1991) permet d'apporter cet aspect supervisé. En effet, quand l'ACP maximise la variance des données, SIR maximise une cohérence entre l'expression du paramètre et la projection des données sur les premières directions renvoyées par la méthode (cf. chap 2.). Nous nous emploierons dans ce chapitre à exécuter la même étude préliminaire que celle utilisant l'ACP au chapitre précédent.

1. L'entropie de Shannon correspond à la quantité d'information qu'il faut stocker ou transmettre

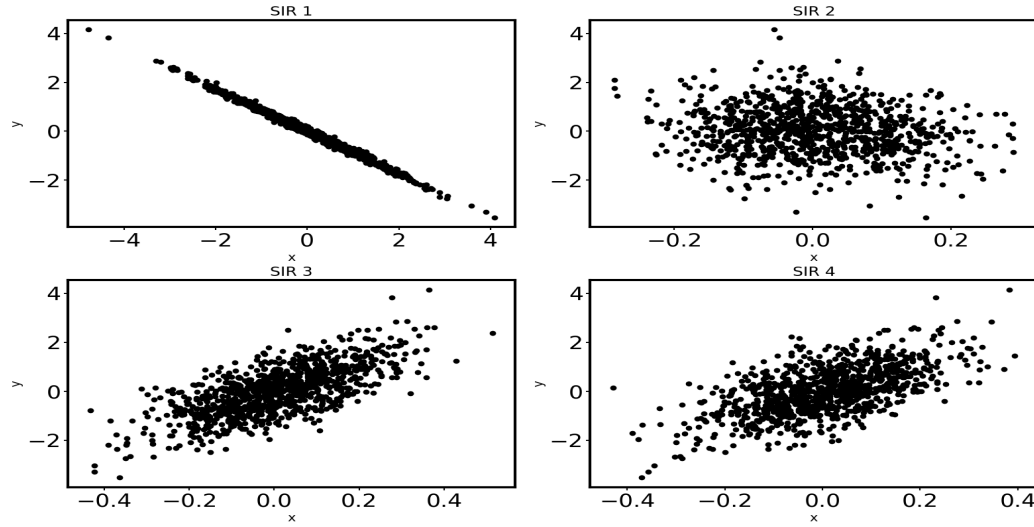


FIGURE 4.1 – Expression de y en fonction des valeurs des individus projetés sur les différentes composantes de SIR (cas de l'exemple linéaire).

4.1 Cas d'un problème linéaire

Ici nous appliquerons la méthode SIR (Li, 1991) (décrite au chapitre 2.) sur le même exemple que pour le paragraphe 3.3.1 du chapitre précédent.

La figure 4.1 est à comparer à celle obtenue avec l'ACP sur le même exemple. On ne peut pas voir de différences qualitatives avec les résultats obtenus avec l'ACP (cf. figure 3.6). On observe, comme dans le cas de l'ACP, une corrélation entre la première composante et y (la variable à expliquer), mais il faut quantifier ces résultats pour savoir si la composante est plus ou moins corrélée que dans le cas de l'ACP. C'est en appliquant les méthodes d'estimation, à partir du paragraphe 4.1.1, que nous pourrions quantifier une différence. De même que dans le chapitre 3, nous appliquerons les méthodes d'estimation basées sur les régressions locales et les k -plus proches voisins.

Nombre de composantes	1	2	3	Pas de SIR
Erreur	0.07	0.063	0.062	0.072

TABLE 4.1 – Table des erreurs moyennes en valeur absolue suivant le nombre de composantes retenues en employant les régressions locales (exemple linéaire). $\sigma_\epsilon = 0.1$.

Nombre de composantes	1	2	3	Pas de SIR
Erreur	0.43	0.44	0.44	0.45

TABLE 4.2 – Table des erreurs moyennes en valeur absolue suivant le nombre de composantes retenues en employant les régressions locales (exemple linéaire) $\sigma_\epsilon = 0.8$.

4.1.1 Régressions locales

Sur l'exemple peu bruité ($\sigma_\epsilon = 0.1$, table 4.1), on obtient des résultats très similaires à ceux obtenus avec l'ACP. La variance des données étant principalement portée par l'information qui nous intéresse, SIR et l'ACP donnent des directions semblables.

Dans le cas bruité, SIR présente un plus gros avantage. Là où précédemment l'ACP dégradait l'erreur d'estimation, SIR parvient à mettre en évidence une direction pertinente même quand celle-ci n'est pas celle qui porte la plus grande variance, c'est-à-dire quand la puissance du signal est faible par rapport à la puissance du bruit (table 4.2).

La table 4.2 montre que même avec de moins bons résultats sur les données brutes (colonne 4)², on obtient systématiquement (quel que soit le nombre de directions retenues) de meilleurs résultats que lors de l'emploi de l'ACP (erreur entre 0.44 sans ACP augmentant jusqu'à 0.49 lors de la conservation de la première composante). La colonne 1 montre que la première composante contient toute l'information pertinente³, car l'ajout d'autres directions pour l'estimation ne fait qu'augmenter l'erreur d'estimation.

2. La différence de résultat sur les données brutes est liée à la réalisation du bruit généré aléatoirement de façon indépendante pour chaque essai.

3. A défaut de contenir toute l'information pertinente, les autres composantes apportent plus de bruit que de signal en terme de puissance.

Nombre de composantes	1	2	3	Pas de SIR
Erreur	0.070	0.075	0.080	0.079

TABLE 4.3 – Table des erreurs moyennes en valeur absolue suivant le nombre de composantes retenues en employant les k-plus proches voisins (exemple linéaire). $\sigma_\epsilon = 0.1$.

4.1.2 k-plus proches voisins

De même qu’au chapitre précédent, on applique une méthode des k-plus proches voisins pour obtenir les résultats présentés table 4.3.

La table 4.3 montre aussi de meilleurs résultats obtenus en employant SIR que lorsque l’on avait employé l’ACP, comme présenté au chapitre précédent. L’ACP affichait alors une erreur de 0.072 à son minimum (alors que la situation était plus favorable étant donné que l’estimation sans projection donnait déjà de meilleurs résultats).

4.1.3 Conclusion sur l’exemple linéaire

Dans cet exemple, on peut voir que SIR présente surtout un intérêt à fort niveau de bruit. Cela est aussi en partie dû à la matrice de variance-covariance, qui, par la structure des données, est mal conditionnée. Il se trouve que le conditionnement de celle-ci diminue quand le bruit augmente. Sur des données plus complexes, si le conditionnement de la matrice de variance-covariance pose problème, nous appliquerons une régularisation.

4.2 Cas d’un problème non-linéaire

Dans ce paragraphe nous confronterons SIR à l’exemple présenté chapitre 3 au paragraphe 3.3.2. Les données $\mathbf{X}$ sont dans un espace de dimension 4, et sont générées à partir des valeurs de paramètre y comme suit :

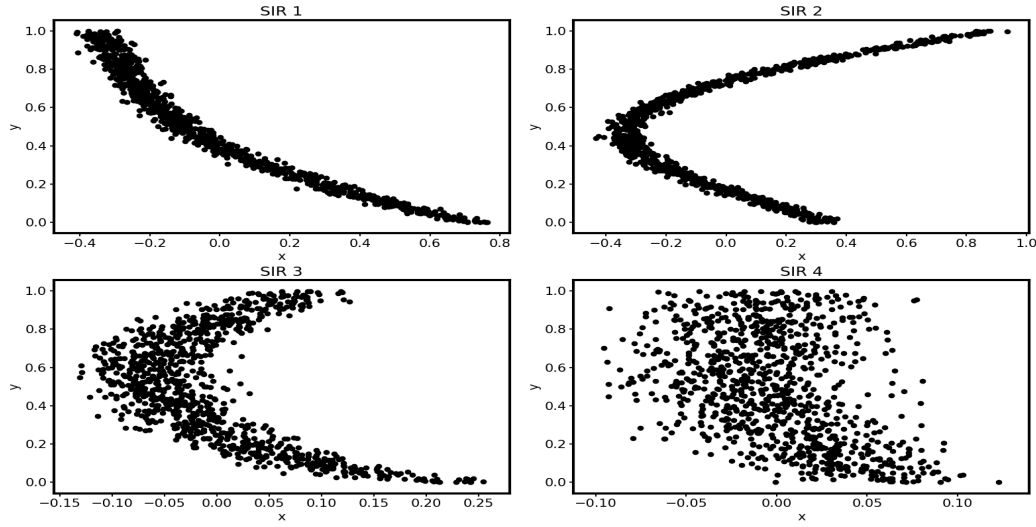


FIGURE 4.2 – Expression de y en fonction des valeurs des individus projetés sur les différentes composantes de SIR (cas de l'exemple non-linéaire).

$$x_1 = y^3 + \epsilon_1 \quad (4.1)$$

$$x_2 = \sin(3y) + \epsilon_2 \quad (4.2)$$

$$x_3 = \sqrt{3} + \epsilon_3 \quad (4.3)$$

$$x_4 = \epsilon_4 \quad (4.4)$$

La figure 4.2 montre que l'expression des individus sur la première direction donnée par SIR, en haut à gauche, présente une corrélation plus "linéaire" (on peut, en effet, approcher le nuage de points par une droite, ce qui est impossible pour ce qui concerne les trois autres directions) que celle obtenue sur la première composante de l'ACP (cf. chapitre 3 paragraphe 3.3.2).

4.2.1 Régressions locales

En appliquant la méthode des régressions locales présentée au chapitre précédent, on obtient les résultats de la table 4.4.

Nombre de composantes	1	2	3	Pas de SIR
Erreur	0.036	0.022	0.025	0.070

TABLE 4.4 – Table des erreurs moyennes en valeur absolue pour l’estimation de y suivant le nombre de composantes retenues en employant les régressions locales (exemple non-linéaire). $\sigma_\epsilon = 0.05$.

Nombre de composantes	1	2	3	Pas de SIR
Erreur	0.025	0.022	0.024	0.025

TABLE 4.5 – Table des erreurs moyennes en valeur absolue pour l’estimation de y suivant le nombre de composantes retenues en employant les k-plus proches voisins (exemple non-linéaire). $\sigma_\epsilon = 0.05$.

Si l’on compare ces résultats à ceux obtenus avec l’ACP, on constate que SIR donne systématiquement de meilleurs résultats que l’ACP. On obtient un minimum dans les deux cas pour deux directions, et l’ACP donne une erreur (0.048) environ deux fois plus élevée que SIR.

4.2.2 k-plus proches voisins

De même, en appliquant la méthode des k-plus proches voisins, on obtient les résultats table 4.5.

La table 4.5 montre que dans le cas de l’emploi de la méthode des k-PPV, SIR permet une réduction de dimension sans perte d’information. Cependant, sur cet exemple, la réduction de dimension ne permet pas de réduire l’erreur moyenne d’estimation de y . Cela sera plus pertinent sur des exemples en dimension plus élevés.

4.2.3 Sélection de directions pertinentes

De même qu’au chapitre 3, nous souhaitons savoir si, sur cet exemple, il ne serait pas pertinent de composer à partir des directions données par SIR un sous-espace individualisé

Réalisation	Avec sélection	Sans sélection	Direction 1	Direction 2	Directions 1+2
1	0.023	0.023	6	0	94
2	0.028	0.025	10	0	90
3	0.02	0.02	12	8	80
4	0.025	0.028	11	0	89
5	0.025	0.026	10	0	90
6	0.024	0.025	12	12	76
7	0.023	0.023	23	13	64
8	0.025	0.026	19	5	76
9	0.025	0.025	10	0	90
10	0.022	0.022	12	0	88
moyenne	0.024	0.0243	12.5	3.8	83.7

TABLE 4.6 – Comparaison des erreurs obtenues avec ou sans sélection des directions pertinentes sur 10 réalisations avec la méthode utilisant les régressions locales. Les trois colonnes de droite représentent le nombre d’individus par essai.

pour chaque élément (individu) de la base.

La table 4.6, met en évidence que l’on peut se contenter de moins de directions pour caractériser un certain nombre d’individus, sans dégrader la qualité de l’estimation de y . Mais, cette caractéristique de l’approche n’est pas celle qui nous intéresse car nous ne souhaitons pas compresser l’information. En effet seule l’estimation précise des paramètres importe dans cette étude. Nous verrons dans la suite, à partir du paragraphe 4.3, que dans le cas d’un espace des données de dimension plus élevée, cette technique de sélection des directions permet une véritable réduction de l’erreur d’estimation.

La table 4.7 montre que lorsque l’on emploie la méthode des k-PPV on arrive à la même conclusion que dans le cas des régressions locales, c’est-à-dire qu’il n’est pas pertinent de conserver les mêmes directions pour tous les individus traités. L’espace de dimension 4 permet de vérifier que l’hypothèse selon laquelle le sous-espace de projection le plus

Réalisation	Avec sélection	Sans sélection	Direction 1	Direction 2	Directions 1+2
1	0.02	0.02	15	11	74
2	0.025	0.025	10	13	77
3	0.020	0.019	14	13	73
4	0.021	0.021	13	16	71
5	0.017	0.018	14	15	71
6	0.021	0.022	11	14	75
7	0.021	0.022	17	20	63
8	0.019	0.019	20	11	69
9	0.026	0.026	13	8	79
10	0.023	0.023	16	20	64
moyenne	0.0213	0.0215	14.3	15.6	70.1

TABLE 4.7 – Comparaison des erreurs obtenues avec ou sans sélection des directions pertinentes sur 10 réalisations avec la méthode utilisant les k-plus proches voisins. Les trois colonnes de droite représentent le nombre d’individus par essai.

pertinent au sens du paramètre, n'est pas le même pour tous les individus, mais cet espace est de dimension trop faible pour que l'on puisse constater une amélioration statistique des résultats.

4.3 Cas d'un problème complexe

En reprenant l'exemple des raies gaussiennes bruitées, dont on souhaite déterminer largeurs et dépressions, on obtient les résultats présentés table 4.8.

Dans cet exemple le conditionnement de la matrice Σ est tel qu'une régularisation est nécessaire (comme expliqué au chapitre 2, paragraphe 2.2.4). On améliorera le conditionnement en appliquant au préalable une régularisation, par troncature des valeurs singulières, sur les données comme nous l'avons fait dans Watson et al. (2017a). On appliquera une réduction de dimension basée sur une ACP en conservant l'espace de plus grande dimension possible qui garantit une inversion stable de la matrice de variance-covariance. Les régressions basées sur SIR seront appliquées sur ce nouvel espace.

La table 4.8 montre que l'erreur est systématiquement inférieure lorsque l'on utilise la sélection des directions quel que soit le niveau de bruit. Cela montre que l'hypothèse, justifiant la sélection des directions, prend un intérêt pratique sur cet exemple plus complexe, dans un espace de plus grande dimension, car cette sélection permet une réduction systématique de l'erreur d'estimation.

L'approche par k-PPV donne des erreurs beaucoup plus faibles (surtout pour le paramètre non-linéaire) que l'approche par régressions locales (cette dernière est intéressante à haut RSB pour le paramètre linéaire). Ainsi, la sélection des directions permet d'optimiser le sous espace de projection indépendamment pour chaque individu, et l'approche par k-PPV est robuste par rapport aux non-linéarités.

4.4 Application sur des données synthétiques

Dans un premier temps, nous appliquerons les différentes méthodes sur les mêmes spectres synthétiques utilisés au chapitre 3. Nous estimerons les valeurs des paramètres de

	k-PPV	k-PPV + sélection	Régression locale	Régression locale + sélection
RSB = 25dB				
Dépressions	0.0048	0.0043	0.0038	0.0020
Largeurs	0.019	0.016	0.049	0.025
RSB = 11dB				
Dépressions	0.015	0.013	0.033	0.016
Largeurs	0.052	0.042	0.080	0.075
RSB = 6dB				
Dépressions	0.026	0.025	0.045	0.026
Largeurs	0.081	0.078	0.11	0.11

TABLE 4.8 – Erreurs moyennes en valeurs absolues obtenues pour l’estimation des dépressions et largeurs des gaussiennes pour différents niveaux de bruit et avec les différentes méthodes.

température effective T_{eff} , de gravité de surface $\log(g)$ de métallicité $[Fe/H]$ et de vitesse de rotation projetée $v \sin(i)$ et pourrons ainsi comparer les résultats des différentes approches basées sur SIR (k-PPV avec et sans sélection). Nous comparerons les résultats obtenus, à titre de référence, avec ceux du chapitre précédent lorsque la projection de faisait sur la base de l’ACP⁴.

La table 4.9 montre que SIR permet une très nette amélioration des résultats pour un haut rapport signal sur bruit. Cependant la sélection de direction ne présente pas d’intérêt dans ce cas. En effet, lorsqu’il n’y a pas de bruit, le voisinage servant à l’estimation est souvent composé des mêmes individus que dans le sous-espace classique obtenu avec SIR.

Si l’on augmente la puissance du bruit, on obtient les résultats présentés table 4.10. Alors, la sélection des directions prend son sens, faisant décroître l’erreur obtenue avec SIR.

4. Les spectres couvrent le domaine spectral de 390 nm à 680 nm avec des valeurs de température allant de 4000 K à 7500 K, des valeurs de gravité de surface allant de 3 à 5 dex des valeurs de métallicité allant de -1 à 1.5 dex et des valeurs de vitesse de rotation projetée allant de 0 à 150 km/s.

	ACP	SIR	SIR avec sélection
Err(T_{eff})(K)	169	30	30
Err(log(g))(dex)	0.3	0.066	0.066
Err([Fe/H])(dex)	0.095	0.0093	0.0095
Err($v \sin(i)$)km/s	1.3	0.086	0.090

TABLE 4.9 – Tableau comparatif des erreurs moyennes d’estimation des 4 paramètres pour un rapport signal sur bruit (RSB) de 45 dB.

	ACP	SIR	SIR avec sélection
Err(T_{eff})	215	83	73
Err(log(g))	0.37	0.13	0.12
Err([Fe/H])	0.12	0.046	0.045
Err($v \sin(i)$)	1.35	0.90	0.34

TABLE 4.10 – Tableau comparatif des erreurs moyennes d’estimation des 4 paramètres pour un RSB de 25 dB.

L’effet est surtout important pour ce qui concerne la température effective T_{eff} et la vitesse de rotation projetée $v \sin(i)$.

On peut conclure de cette première application que l’emploi de SIR, en lieu et place de l’ACP, pour la réduction de dimensionnalité visant à la facilitation de l’estimation des paramètres stellaires fondamentaux, permet de réduire de manière significative l’erreur d’estimation. L’approche originale proposée consistant à sélectionner des directions particulières formant un sous-espace différencié pour chaque individu, ne présente en fait un réel intérêt qu’à faible rapport signal sur bruit (comparable à celui qu’on attend sur les données réelles).

	ACP	SIR	SIR avec sélection
Err(T_{eff})	102	88	84
Err($\log(g)$)	0.20	0.17	0.17
Err($[\text{Fe}/\text{H}]$)	0.087	0.068	0.068
Err($v \sin(i)$)	1.37	1.38	1.38

TABLE 4.11 – Tableau comparatif des erreurs moyennes d’estimation des 4 paramètres pour S4N à partir d’Elodie.

4.5 Application aux données réelles

Ici nous confronterons les approches basées sur SIR, présentées dans ce chapitre, aux données réelles, utilisées dans Paletou et al. (2015a). Ces données sont constituées de 905 spectres, observés par l’instrument Elodie pour la base de données de référence, et de 104 spectres issus de la campagne de mesure du S4N pour la base de test.

La table 4.11 montre l’étude comparative des résultats obtenus avec l’ACP et SIR pour l’estimation des paramètres des 104 individus du S4N sélectionnés à partir des 905 spectres de la base Elodie. Hormis pour ce qui concerne la vitesse de rotation projetée SIR permet de réduire l’erreur sur les paramètres estimés. Ici, il n’y a pas non plus intérêt à sélectionner les directions pertinentes, car celles-ci donnent lieu au même voisinage que l’espace donné par SIR. Le fait que SIR ne permette pas d’améliorer les résultats concernant la vitesse de rotation projetée, est sûrement liée à l’échantillonnage des vitesses de la base Elodie. En effet, il y a peu de spectres à haute vitesse et beaucoup plus de spectres à faible vitesse et cela a un impact lors de la construction des tranches qui ne permet pas à SIR de fonctionner dans des conditions optimales. L’échantillonnage des autres paramètres, sans être régulier, est plus homogène, ce qui explique que les résultats avec SIR, bien que toujours meilleurs que ceux obtenus avec l’ACP, présentent une amélioration moins nette que dans le cas des données synthétiques. Dans le cas des données synthétiques les trois paramètres fondamentaux étaient échantillonnés de manière uniforme, et la vitesse de rotation projetée

était échantillonnée de manière non-uniforme, mais de façon beaucoup plus dense que dans le cas des données réelles⁵. Sur les données réelles SIR permet d’avoir des résultats plus performants que l’ACP, mais de manière moins marquée. De la même manière, la sélection des directions pertinentes n’est pas efficace comme pour les données synthétiques.

4.6 Conclusion sur l’utilisation de SIR

L’emploi de SIR pour l’estimation de paramètres stellaires fondamentaux présente un intérêt démontré par les résultats présentés précédemment. En effet, cette méthode permet, à l’instar de l’ACP, de traiter les données sans avoir nécessairement connaissance de leur nature. Il n’est jamais nécessaire d’apporter une éventuelle connaissance d’un modèle mathématique ou d’un lien fonctionnel quelconque entre les données et les paramètres. Par contre, par rapport à l’ACP, SIR prend en compte la connaissance des valeurs de paramètres de la base de référence au moment de la construction du sous-espace. L’emploi de SIR permet donc d’obtenir de meilleurs résultats qu’avec l’ACP et ce malgré la difficulté posée par le conditionnement de la matrice de variance-covariance Σ . La sélection des directions, pertinentes quant à elles, présente surtout un intérêt à faible rapport signal sur bruit. L’hypothèse de départ est que l’information qui nous intéresse n’est pas toujours portée par les mêmes directions. Les premières directions peuvent dans certaines zones du sous-espace être davantage porteuses de bruit et il est attendu que cette hypothèse se vérifie mieux lorsque le bruit est puissant. L’approche par sélection des directions peut aussi avoir un intérêt périphérique à cette étude. Il est effectivement possible de l’employer dans un contexte de compression de l’information. On peut, grâce à cette sélection, avoir une information tout aussi précise en conservant moins de données.

L’étude comparative de l’ACP et de SIR pour la détermination de paramètres physiques à partir d’une base de données de référence a, par ailleurs, donné lieu aux publications qui suivent aux conférences ECMSM (Watson et al., 2017a) et GretsI (Watson et al., 2017b).

5. Les valeurs allant de 0 km/s à 150 km/s avec un pas d’échantillonnage croissant, mais avec le même nombre de spectres par valeur de vitesse.

Inférence of an explanatory variable from observations in a high-dimensional space : Application to high-resolution spectra of stars

Résumé : Notre but est de déterminer les paramètres fondamentaux à partir de l'analyse du spectre électromagnétique d'une étoile. Nous utiliserons entre 10^3 et 10^5 spectres, chacun d'entre eux étant un vecteur de données de dimension 10^2 à 10^4 . Nous sommes ainsi confrontés à des problèmes en lien avec la "malédiction de la dimensionnalité". Nous recherchons une méthode qui réduise la taille de cet espace des données en ne conservant que la partie la plus pertinente de l'information. Nous utiliserons comme méthode de référence l'analyse en composantes principales (ACP). Cependant l'ACP est une méthode non-supervisée, ainsi le sous-espace qu'elle propose n'est pas cohérent avec les variations du paramètre. Nous avons donc testé une méthode supervisée basée sur la régression inverse par tranches (SIR) qui doit fournir un sous-espace cohérent avec le paramètre étudié. On peut voir des analogies avec l'analyse factorielle discriminante dans un contexte de classification : la méthode découpe la base de données en tranches suivant la variation du paramètre étudié, et compose un sous-espace qui maximise la variance inter-tranches tout en normalisant la variance totale des données projetées. Néanmoins, les performances de SIR ne sont pas toujours satisfaisantes dans son usage standard et cela est dû à la non-monotonie de la fonction inconnue liant le paramètre aux données. Nous montrons que dans certains cas on peut obtenir de meilleurs résultats en sélectionnant les directions les plus pertinentes pour la détermination de la valeur du paramètre. Des tests préliminaires sont effectués sur des pseudo-raies synthétiques bruitées. On peut voir que lorsque l'on ne retient que la première direction, SIR permet de faire décroître l'erreur d'estimation de 50% pour ce qui concerne le paramètre non-linéaire et de 70% pour ce qui concerne le paramètre linéaire. De plus, lorsque l'on sélectionne de manière différenciée la direction la plus pertinente, l'erreur décroît de 80% pour le paramètre non-linéaire et de 95% pour le paramètre linéaire.

Inference of an explanatory variable from observations in a high-dimensional space: Application to high-resolution spectra of stars

Victor Watson^{1,2}, Jean-François Trouilhet^{1,2} (IEEE senior member), Frédéric Paletou^{1,2} and Stéphane Girard³

(1) Université de Toulouse, UPS-Observatoire Midi-Pyrénées, Irap, Toulouse, France

(2) CNRS, Institut de Recherche en Astrophysique et Planétologie, 14 av. E. Belin, 31400 Toulouse, France

(3) INRIA, 655 Avenue de l'Europe, 38330 Montbonnot-Saint-Martin, France

Email : {victor.watson, jean-francois.trouilhet, frederic.paletou}@irap.omp.eu , stephane.girard@inria.fr

Abstract—Our aim is to evaluate fundamental parameters from the analysis of the electromagnetic spectra of stars. We may use 10^3 - 10^5 spectra; each spectrum being a vector with 10^2 - 10^4 coordinates. We thus face the so-called “curse of dimensionality”. We look for a method to reduce the size of this data-space, keeping only the most relevant information. As a reference method, we use principal component analysis (PCA) to reduce dimensionality. However, PCA is an unsupervised method, therefore its subspace was not consistent with the parameter. We thus tested a supervised method based on Sliced Inverse Regression (SIR), which provides a subspace consistent with the parameter. It also shares analogies with factorial discriminant analysis: the method slices the database along the parameter variation, and builds the subspace which maximizes the inter-slice variance, while standardizing the total projected variance of the data. Nevertheless the performances of SIR were not satisfying in standard usage, because of the non-monotonicity of the unknown function linking the data to the parameter and because of the noise propagation. We show that better performances can be achieved by selecting the most relevant directions for parameter inference.

Preliminary tests are performed on synthetic pseudo-line profiles plus noise. Using one direction, we show that compared to PCA, the error associated with SIR is 50% smaller on a non-linear parameter, and 70% smaller on a linear parameter. Moreover, using a selected direction, the error is 80% smaller for a non-linear parameter, and 95% smaller for a linear parameter.

I. INTRODUCTION

As an historically data-intensive science, astrophysics has obviously to deal with significant data-processing issues. Projects such as the Gaia survey [1] or LSST¹ are acquiring and will acquire a lot of data, both in terms of number of samples and in terms of amount of data by sample. We are interested in methods able to extract as much physical content as possible from such large volumes of data.

Our problem deals with high dimensional data-processing. Similarly to [2], our goal is to infer stellar fundamental parameters such as effective temperature, surface gravity or metallicity. We also want to infer an important physical parameter which may significantly affect the appearance of stellar spectra, the projected rotational velocity.

These physical parameters values are to be estimated from the stars measured spectra. Nowadays precision optical spectroscopy relies on observations made over large spectral bandwidths, of the order of about 500 nm or more (without any gap), and at high-spectral resolution, typically of several 10^4 in resolvance (resolvance is related to sampling properties of these instruments $R = \frac{\lambda}{\Delta\lambda}$).

For these problems, there is no direct linear relationship between the data and the physical parameters. To ease the search for such

a function or relationship, we use methods based on a reduction of the dimensionality of our data-space. These methods should extract from the data the information related to the physical parameter we want to infer. One of the most popular among these methods is PCA [3] which is the method used by [2]. In this study, our aim is to find the best subspace for the estimation of an explanatory variable y from a data vector x . This subspace is to be estimated from a database containing many examples of pairs (x_i, y_i) . Once the subspace is found, a k-Nearest-Neighbours method will give the estimator $\hat{y}_i$ of the unknown value y_i associated with the data vector x_i . Hereafter we will propose an alternative solution based on SIR [4] to improve the results. Unlike PCA, SIR is a supervised method. Hence, when building the projection subspace, we will be able to add the information about the physical parameters (or explanatory variables) we want to retrieve from the data. These parameters are considered independently one at a time. SIR method has already been used in [5]. One of the differences with our case is that we have no information about the dimension of the subspace consistent with the one-dimensional manifold described by the physical parameter; and we cannot use a physical modelling to determine the dimension of this subspace.

SIR suffers from a limitation due to the ill-conditioning of the variance-covariance matrix. Thus, it will sometimes require some preprocessing in order to perform a stable inversion of the variance-covariance matrix. This preprocessing will consist in applying a PCA on the data before applying SIR.

The first section briefly recalls the basic principles of PCA and SIR methods. The second section presents the implementation of SIR and some examples to enlighten how it basically works. The third section is about testing SIR method on a simplified data set which mimics the astrophysical data we want to deal with. The fourth section discusses the tuning of the method and the pre-processing.

II. METHODS

A. Principal Component Analysis

PCA is a method which aims at finding the subspace that preserves the most the variance of the projected data.

Considering a data matrix $\mathbf{X}$ of dimension $M \times N$ with each line being an element from the database, PCA finds the projection matrix $\mathbf{V}$ such as $\mathbf{X}^T \mathbf{V}$ has the maximum variance. The first vector v_1 is the one that maximizes $Var(\mathbf{X}^T v_1)$ under the constraint $v_1^T v_1 = 1$, the second vector v_2 is the vector orthogonal to v_1 that

¹<https://www.lsst.org/>

maximizes $Var(\mathbf{X}^T v_2)$ under the same constraint, and so on until the number of components reaches the specified dimension of the projection subspace. To find these vectors, one can perform an eigen-decomposition of the variance-covariance matrix of $\mathbf{X}$ denoted by Σ . Thus we will find these vectors by solving: $(\Sigma - \lambda \mathbf{I}_M)v = 0$ [3].

The first step of PCA is to compute the variance-covariance matrix Σ :

$$\Sigma_{i,j} = \frac{1}{M} \sum_{k=1}^M (\mathbf{X}_{k,i} - \bar{\mathbf{X}}_i)(\mathbf{X}_{k,j} - \bar{\mathbf{X}}_j) \quad (1)$$

where N is the dimension of the original data-space, $\mathbf{X}_k$ is the k^{th} column containing all the x vectors projected on the k^{th} component of the original data-space. $\bar{\mathbf{X}}_i$ and $\bar{\mathbf{X}}_j$ are the mean values computed on lines i and j .

The subspace obtained by PCA is spanned by the eigenvectors of Σ associated to its greatest eigenvalues.

Let $\mathbf{V}$ be the $N \times P$ matrix with every column being one of these P selected eigenvectors, and N being the dimension of the original space. The x vectors can be projected on the subspace associated with PCA:

$$x_p = (x^T - \bar{\mathbf{X}})\mathbf{V} \quad (2)$$

where x is the representation of an element from the database in the original data-space, and $\bar{\mathbf{X}}$ is the mean of the elements from the database in the original data-space. $\bar{\mathbf{X}}$ is computed on all the lines of the matrix $\mathbf{X}$.

The main drawback of PCA on our application is that this method does not take into account the prior knowledge on the explanatory variable we want to infer, as it is an unsupervised method. This is why we decided to investigate SIR as an alternative.

B. Sliced Inverse Regression

SIR [4] aims at finding the directions that explain the best the variation of the explanatory variable y (taken one by one) from the data x . The principle of the method is very close to the one of the PCA: it relies on the estimation of a linear projection subspace. However SIR takes into account the information on the variation of the explanatory variable in the building of this subspace. The database is divided into H slices along the variation of the explanatory variable, each slice being composed of pairs (x,y) which have close values of the explanatory variable. SIR builds a subspace that maximizes the variance between the slices while minimizing the variance within the slices. Somehow, SIR can be compared to a Linear Discriminant Analysis (LDA)[6] in a classification problem. Thus, SIR finds a linear projection that drives apart the elements that have very different values of the explanatory variable while bringing together the elements which have close values of the explanatory variable. Let us denote by $\mathbf{X}_h$ the elements from the database in the slice h . For every slice we compute the mean $\bar{m}_h$

$$\bar{m}_h = \frac{1}{n_h} \sum_{\mathbf{X}_i \in h} \mathbf{X}_i \quad (3)$$

where n_h is the number of elements in the slice h , so that we have the matrix $\mathbf{X}_H$ with every line being the mean $\bar{m}_h$ associated with the elements of the slice h . We then compute the variance-covariance of the slices means:

$$\Gamma = \frac{1}{H} (\mathbf{X}_H - \bar{\mathbf{X}}_H)^T (\mathbf{X}_H - \bar{\mathbf{X}}_H). \quad (4)$$

The directions computed by SIR will be given by the eigenvectors associated with the greatest eigenvalues of $\Sigma^{-1}\Gamma$. The vectors spanning the sub-space we are looking for are the solutions of $(\Sigma^{-1}\Gamma - \lambda \mathbf{I}_M)v = 0$.

III. SLICED INVERSE REGRESSION

A. Implementation

SIR can be computed the following way:

- compute the variance-covariance matrix Σ as in eq. 1
- sort the parameter vector y so it becomes y_{sorted}
- split the sorted vector y_{sorted} in H non-overlapping slices
- for each slice h compute the mean $\bar{m}_h$ according to eq. 3
- set the Γ matrix as in eq. 4
- build a projection matrix $\mathbf{V}$ so that the columns of $\mathbf{V}$ are the eigenvectors associated with the greatest eigenvalues of $\Sigma^{-1}\Gamma$.

One of the critical point of this method may be the inversion of Σ . Indeed, Σ can be ill-conditioned as discussed in section 4. In that case, we propose to apply a PCA to project the data on a subspace that maximizes the variance. Doing so, PCA will lead to an inversion of Σ more stable, and in the meantime, it will reduce the dimension of the dataspace and thus requires less computational resources. There is a compromise to find between improving enough the conditioning and loosing information held by the data. This will be discussed in section 5.

B. Examples

1) *Linear example:* As a first example we will consider a linear relationship between x and y such as $y = \beta^T x + \epsilon$. The dimension of x is 4 and the database contains 1000 elements. $\mathbf{X}$ follows a normal standard distribution $\mathbf{X} \sim N(\mu, I_4)$, where μ is a null vector of dimension 4×1 and I_4 is a dimension 4 identity matrix, $\beta^T = [1, 1, 0, 1]$ and the noise $\epsilon \sim N(0, 0.1I_4)$.

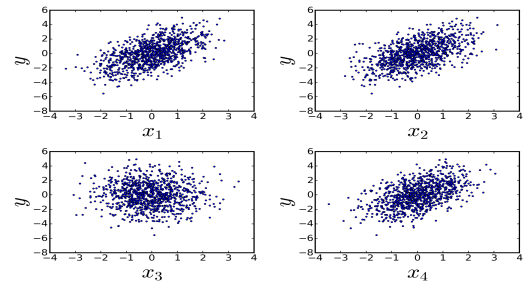


Fig. 1. Values of y versus the realisations of $\mathbf{X}$ for the example of section III-B1. We can observe that none of these directions x_i of the original data-space is relevant to explain the variations of the explanatory variable y . This can be explained by the x_i varying independently from one to another, y being a linear combination of these x_i , and all these components having the same weight.

Fig. 1 shows that from every component of $\mathbf{X}$, it is impossible to determine precisely y . It appears in these four sub-figures that for any value of the x_i , y looks like it is randomly distributed. This is why one must find the correct linear combination of the x_i to explain y .

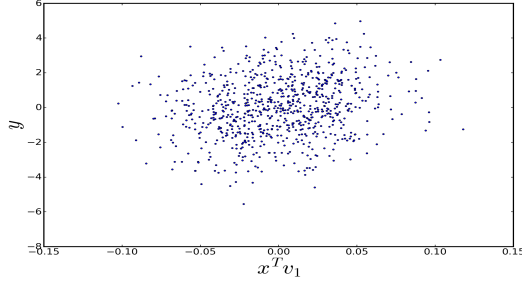


Fig. 2. Values of y versus the realisations of $\mathbf{X}$ projected on the first direction given by PCA for the example of section III-B1. Such a cloud of points indicates that the variance is nearly irrelevant to infer the value of y from a known value of x .

The first direction v_1 given by PCA, displayed in Fig. 2, is the one that maximizes the variance in the data, and it appears that this criterion is not relevant here, because as in Fig. 1, y seems to be randomly distributed for any value of the projected data $x^T v_1$.

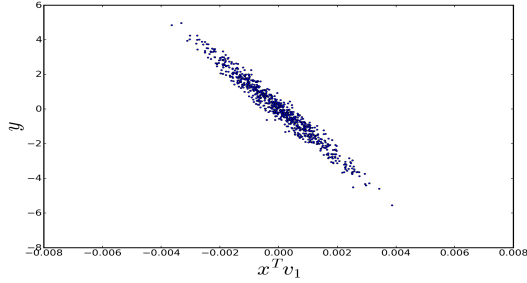


Fig. 3. Values of y versus the realisations of $\mathbf{X}$ projected on the first direction given by SIR for the example of section III-B1. This figure clearly shows that SIR is able to find the linear projection which explains y from the components presented on Fig.1.

Fig. 3 shows that SIR provides a direction v_1 that allows us to determine quite precisely y from $x^T v_1$. Indeed for every value of $x^T v_1$ the range of values that y can take is limited.

This first example shows that the first direction given by SIR (Fig. 3) finds a much more relevant combination of $\mathbf{X}$ to explain y than PCA (Fig. 2) even though it is quite difficult to infer that relationship from $\mathbf{X}$ in the original space (Fig. 1). In this case, the condition number associated with Σ is 1.26, so there is no need for any preprocessing.

2) *Nonlinear example:* The second example focusses on a non linear case. All the components of $\mathbf{X}$ are independent and follow a standard uniform distribution $U[0, 1]$ and $y(x) = 2x_1x_2 + x_3 + 0x_4 + \epsilon$, where the noise is $\epsilon \sim N(0, 0.01I_4)$.

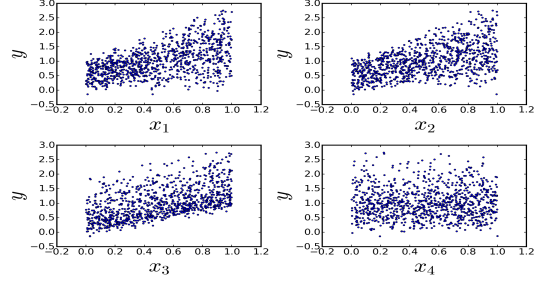


Fig. 4. Values of y versus the realisations of $\mathbf{X}$ for the example of section III-B2. As in the case of the linear example, it is not obvious to find a relationship between the data in its original space and the explanatory variable even if we can see that the three first components (top-left, top-right, and bottom-left) of the data-space hold information about y . This is due to the fact that the range of the values taken by y is rather wide whatever the value of x_i is.

Once again in this example, we have to find a good combination of the components x_i shown on Fig. 4 to explain y .

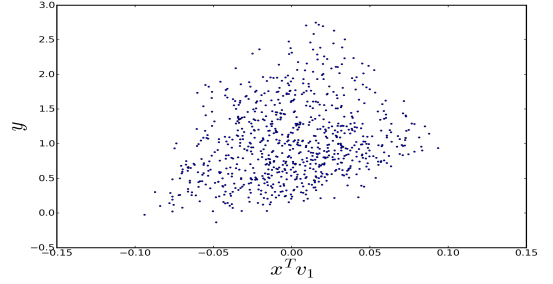


Fig. 5. Values of y versus the realisations of $\mathbf{X}$ projected on the first directions given by PCA for the example of section III-B2. It appears that the first direction given by PCA holds almost no information about the explanatory variable because no matter what $x^T v_1$ is, y cannot be precisely determined.

We show with Fig. 5 that, using PCA, a relationship appears between $x^T v_1$ and y . Indeed, the scatter-plot (Fig. 5) exhibits a shape and is not fully random. Nevertheless, this so-called 'shape' is far from being sufficiently narrow to determine precisely y from $x^T v_1$.

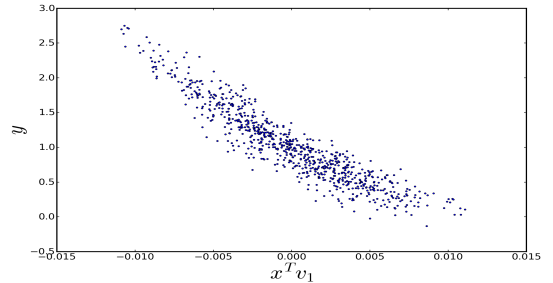


Fig. 6. Values of y versus the realisations of $\mathbf{X}$ projected on the first direction given by SIR for the example of section III-B2. Even though the direction is not as efficient as in the linear case, we can see here that SIR projects the data on a direction which highly correlates the data with the explanatory variable.

Again in the case where the relationship between the data and the parameter is not linear, SIR (Fig. 6) provides more relevant directions

to explain y from $\mathbf{X}$ than PCA. Of course, since it is a linear projection method more than one direction can be necessary to estimate correctly the one dimensional explanatory variable which is non-linearly linked to the data. In this example, a preprocessing of the data was not necessary because the condition number associated with Σ is 1.17.

IV. APPLICATION TO SYNTHETIC DATA

A. Inference of simple Gaussian lines parameters

We showed that SIR can determine the combinations of $\mathbf{X}$ that best explain y . We will try to use it to obtain $\hat{y}$ when it is unknown. Hereafter, the database is a set of simplified Gaussian lines that mimic the astrophysical data of interest. It is a set of 10^3 Gaussian $\mathbf{X}$ in which 800 will be used as a reference (y is known) database $\mathbf{X}_0$ and the 200 others will be used as a test database $\mathbf{X}_i$ (y is to be determined). Each Gaussian x_j is associated with a set of explanatory variables h_j and w_j , respectively the depression and the width of a simplified Gaussian spectral line:

$$x(\lambda; h, w) = 1 - h \times e^{-\frac{1}{2} \left(\frac{\lambda - \lambda_0}{w} \right)^2} + \epsilon. \quad (5)$$

h and w will be processed independently, $\lambda_0 = 0$ is common to all the elements. h and w are uniformly distributed $h \sim U[0.1, 0.9]$, $w \sim U[1, 3]$ and the noise ϵ is normally distributed $\epsilon \sim N(0, \sigma_\epsilon)$.

For this example we will add a noise ϵ such as $\sigma_\epsilon = 0.005$ equivalent to a signal to noise ratio around 32dB.

Once the data is projected on the subspace computed by SIR we will proceed with a simple mean on a 10 nearest neighbours method, to infer the value of the explanatory variable $\hat{y}_i$ of an individual x_i :

$$\hat{y}_i = \frac{1}{10} \sum_{y_{0_j} \in W_i} y_{0_j}. \quad (6)$$

Here, W_i is the set of the values of the explanatory variable from the training data-base x_p , associated with the 10 nearest neighbours of x_{p_i} .

Because of the ill-conditioning of Σ (about $3 \cdot 10^5$), we chose to perform a PCA first on the data, in order to improve the conditioning. SIR is ran on the data projected on the subspace of dimension 6 given by PCA, so that we can see the directions given by SIR on Fig. 7 and Fig. 8. The choice of the dimension of the subspace given by PCA as a preprocessing will be discussed in section 5.

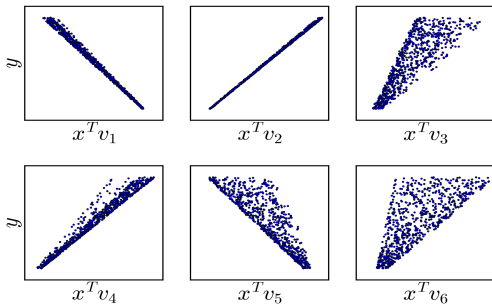


Fig. 7. Values of the depressions of the Gaussians y versus the realisations of $\mathbf{X}_0$ projected on the 6 directions given by SIR. Most of the information we seek is held by the first two directions (top-left and top-center), but in some cases the directions 3 (top-right), 4 (bottom-left) or 5 (bottom-center) can add quite precise information about y .

We can see on Fig. 7 that most of the information about the depression is held by the first 2 directions (top-right and top-center) given by SIR. Indeed we can get an accurate estimator of the depressions from the projection of the data on the second direction (top-center of Fig. 7) as these two quantities are linearly related.

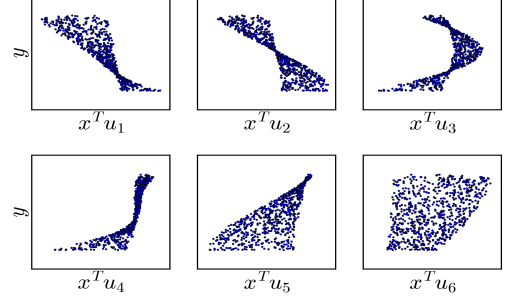


Fig. 8. Values of the widths of the Gaussians y versus the realisations of $\mathbf{X}_0$ projected on the 6 directions given by SIR. In that case the link between the explanatory variable and the data is non-linear so it does not seem to be a subspace which explains correctly all the values that can be taken by y .

We show on Fig. 8, that the direction which is the most relevant to explain the width from the data highly depends on the value of the projected data on the several directions. For instance the first direction is relevant for high values of $x^T u_1$ as the associated values of y are narrow (top-left of Fig. 8). The same first direction though is not very relevant to infer y when the value of $x^T u_1$ is around 0 because the corresponding values of y spread from 1.5 to 3 which is 3/4 of the total range of values y can take.

From Fig. 7 and Fig. 8, one can see that the directions which are relevant for estimating y can vary from an individual to another. So the question remains: which of these directions will be relevant to estimate y from x ? We will suggest a way in the following section to individually select the most relevant sub-space for each of the elements x one wants to determine the paired y .

B. Direction selection

As shown in Fig. 7 and 8, the directions computed by SIR that best explain y from x may depend on the value of x . This is why we suggest to select the directions that will minimize the variance on the values of y associated with the $x^T v$ which are in the neighbourhood of the x_i one wants to determine the associated value $\hat{y}_i$. We thus can have a variable number of directions that is optimal to determine $\hat{y}_i$ from an individual to another. For every individual x_i , we look for the set of directions D_i spanning the subspace which minimizes the variance of $y_0 \in W_i$, where $W_i = \{y_{0_j} | dist_{s_{i_j}} \leq dist_{10_i}\}$, $dist_{i_j} = \|x_i^T D_i - x_j^T D_i\|$ and $dist_{10_i}$ is the distance between $x_i^T D_i$ and its 10^{th} nearest neighbour.

C. Validation of the modified SIR method

The first test we will run is a comparison of the results given with the first direction of PCA, the first direction of SIR, and one selected direction of SIR.

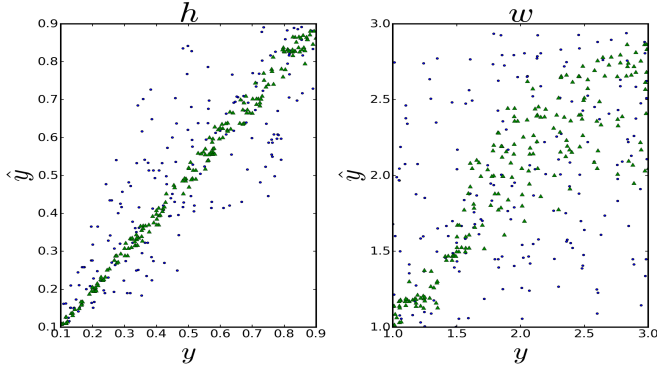


Fig. 9. Results of the estimation of depressions and widths with PCA using only the first direction (blue spots), SIR using the first direction (green triangles).

We can observe in Fig. 9 that the results given by SIR are much more precise than the ones given by PCA. The perfect estimation being the line $\hat{y} = y$, results given by the method using SIR are closer to that line than the PCA ones.

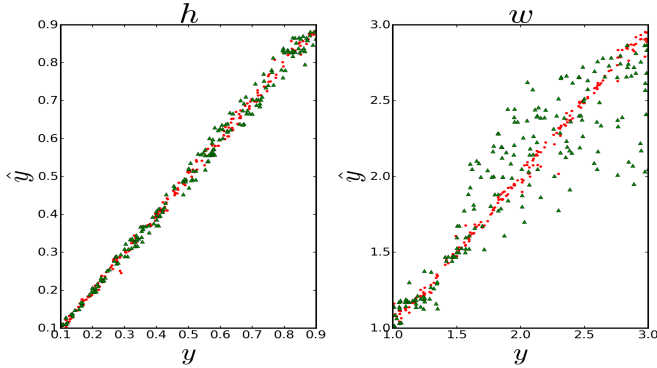


Fig. 10. Results of the estimation of depressions and widths with SIR using the first direction (green triangles), and the modified version of SIR with the direction selection (red spots).

As shown in Fig. 10, the selection of the direction that minimizes the variance of y_0 in the neighbourhood improves the results of SIR, and the efficiency of these different methods can be assessed with the computation of the mean absolute error (MAE) defined as:

$$MAE = \frac{1}{M} \sum_{i=1}^M \|\hat{y}_i - y_i\| \quad (7)$$

	Results with PCA first direction	Results with SIR first direction	Results with SIR one selected direction
Depressions	0.097	0.012	0.010
Widths	0.616	0.223	0.029

TABLE I. Table of the mean absolute errors given by the three versions of the method using only one direction for the projection.

The improvement given by the selection of the directions is very efficient in finding the best direction when the link between x and y is non linear.

We thus can compare the results given by the optimal version of PCA (keeping 6 components) with SIR on its optimal number of directions for either depression (2 directions) or widths (3 directions), and also the version of SIR which looks for the subspace which minimizes the variance on y in the neighbourhood. The results of this comparison are shown on table II.

	PCA	Classic SIR
Depressions	0.0063	0.0025
Widths	0.029	0.027

TABLE II. Table of the mean absolute errors given by PCA with 6 directions, classic SIR with the optimal number of directions, and SIR with the selection of the directions which minimize the variance on y_0 in the neighbourhood of the individual one wants to estimate the value of y .

The results displayed on table II show that SIR provides better results than the ones we can achieve with PCA. The fact that the classical version of SIR yields accurate estimates regarding the depressions can be explained by its efficiency on linear relationships between x and y . Maybe the criterion used to determine the best directions combination (variance on y) is not the best criterion. However there is a non negligible gain regarding the widths estimation with SIR combined with the direction selection. We can note that the dimension of the subspace varies from an individual to another as shown on Fig. 11 for the depressions and Fig. 12 for the widths.

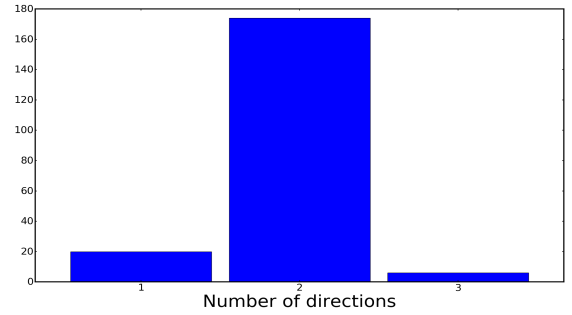


Fig. 11. Histogram of the dimension of the subspaces chosen by the method for the depression estimation for each of the 200 elements of the test database.

We can observe in Fig. 11 that, as in the classical SIR case, most of the time the subspace that best explains the depression value from the data is spanned by 2 directions.

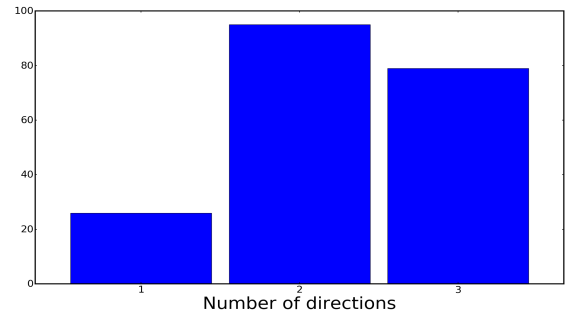


Fig. 12. Histogram of the dimension of the subspaces chosen by the method for the widths estimation for each of the 200 elements of the test database.

Fig. 12 shows that the optimal number of directions to determine the width can either be 2 or 3. With classical SIR the optimal number

of directions is 3, but Fig. 12 shows that it is not that often the best solution.

V. DISCUSSION ABOUT PREPROCESSING

We chose to perform a PCA to improve the conditioning of Σ , as it is a way to reduce the dimension of the data-space, while loosing the least information. Doing that, we assume that the information held by the last directions given by PCA are, even if relevant, drawn into noise. Performing a PCA carefully should not cause any loss of information, except for the information we could not have retrieved anyway. But how many directions should be kept? If we keep too many, we will not make the conditioning decrease enough, and the directions given by SIR will not be accurate. If we do not keep enough directions, there is a risk of loosing information on the explanatory variable.

To get an idea of what the best solution would be, let us have a look at Fig. 13 and Fig. 14 to the 10 first directions given by PCA on our dataset:

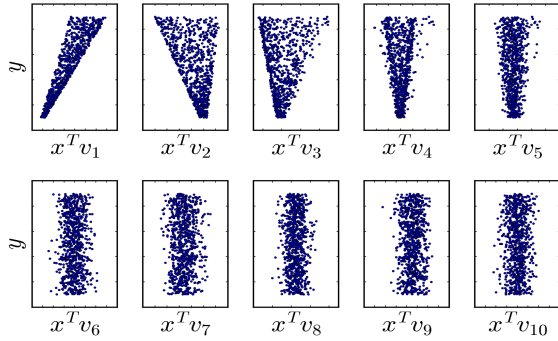


Fig. 13. Value of the depressions of the Gaussians y versus the realisations of $\mathbf{X}_0$ projected on the 10 directions given by PCA. Directions given by PCA seem to enclose information about y up to the fifth direction.

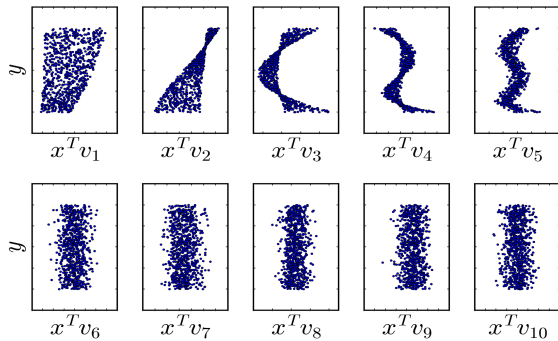


Fig. 14. Value of the widths of the Gaussians y versus the realisations of $\mathbf{X}_0$ projected on the 10 directions given by PCA. As in Fig. 13, the directions over the fifth seems to be irrelevant as the distribution of y looks independent from the value of $x^T v_i$.

We can see in Fig. 13 and Fig. 14 that from the sixth direction (bottom-left), the remaining directions do not seem to hold any information about the parameter, since the value of y looks independent from $x^T v_i$. This assumption is verified by a validation

protocol based on the estimation of the explanatory variables with PCA for several noise realisations. It appears that the dimension of the subspace given by PCA that gives the best results is 6. It makes the conditioning decrease to about $4 \cdot 10^4$. Even if we only gain a 10 factor on the conditioning, we cannot reduce more the dimensionality because we would most probably loose important information about the explanatory variable.

VI. CONCLUSION AND FUTURE WORK

In this article we have shown that SIR, being a supervised method, provides a subspace more relevant to link the data to the explanatory variable than the PCA one. We have also tried to give an answer to the problem of the number of directions which would be relevant. We also considered the problem of the choice of that relevant directions which can change with the location of the studied sample. We have shown that the method can be adapted to give the optimal subspace to locally contain the manifold where the explanatory variable is best explained. Preliminary results for this method using the same data set as in [2], show that, compared to the PCA-based method, the proposed SIR-based method decreases the estimation error by 20% for the effective temperature and the surface gravity while we obtain similar estimation errors for the metallicity. However, the error increases by 30% for the projected rational velocity. The adaptation of the method to real astrophysical data is still in progress. Because the data are a lot more complex, the dimension of the original data-space can be 10 or 100 times larger than the one used for this study, and the conditioning of Σ can reach 10^{23} . We will need to run more tests to determine the optimal number of components to keep with PCA as a preprocessing, and also to know the best way to consider the neighbourhood depending on the number of directions kept by the method. We believe that better results should be achieved by finding a way to correctly tune the method for this particular case.

REFERENCES

- [1] The Gaia Collaboration, “The Gaia mission,” *Astronomy & Astrophysics*, vol. 595, p. A1, 2016.
- [2] F. Paletou, T. Böhm, V. Watson, and J.-F. Trouilhet, “Inversion of stellar fundamental parameters from Espadons and Narval high-resolution spectra,” *Astronomy & Astrophysics*, vol. 573, 2015.
- [3] I. Jolliffe, *Principal Component Analysis*. Springer Verlag, 1986.
- [4] K. C. Li, “Sliced Inverse regression for dimension reduction,” *Journal of the american statistical association*, vol. 86, pp. 316–327, Jun. 1991.
- [5] C. Bernard-Michel, S. Douté, M. Fauvel, L. Gardes, and S. Girard, “Retrieval of Mars surface physical properties from OMEGA hyperspectral images using regularized sliced inverse regression,” *Journal of geophysical research*, vol. 114, 2009.
- [6] G. J. McLachlan, *Discriminant Analysis and Statistical Pattern Recognition Volume 544 of Wiley Series in Probability and Statistics, ISSN 1940-6517 Wiley series in probability and mathematical statistics: Applied probability and statistics Wiley-Interscience paperback series*. John Wiley & Sons, 2004.

Sliced Inverse Regression pour la détermination des paramètres stellaires fondamentaux

Victor WATSON^{1,2}, Jean-François TROUILHET^{1,2}, Frédéric PALETOU^{1,2} Marwan GEBRAN³

¹Université de Toulouse, UPS-Observatoire Midi-Pyrénées, Irap, Toulouse, France

²CNRS, Institut de Recherche en Astrophysique et Planétologie, 14 av. E. Belin, 31400 Toulouse, France

³Department of Physics and Astronomy, Notre Dame University-Louaize, PO Box 72, Zouk Mikael, Lebanon

victor.watson@irap.omp.eu, jean-francois.trouilhet@irap.omp.eu,
frederic.paletou@irap.omp.eu, mgebran@ndu.edu.lb

Résumé – Nous voulons déterminer la valeur d’un paramètre à partir d’un grand vecteur de données associé, sans connaître la fonction permettant de passer d’un espace à l’autre. La méthode *Sliced Inverse Regression* (SIR) permet de projeter un vecteur de données vers un sous-espace qui est cohérent avec la variation du paramètre étudié. Nous proposerons, sur la base du sous-espace obtenu grâce à SIR, une méthode pour estimer de la manière la plus précise possible la valeur que prend le paramètre auquel on s’intéresse.

Abstract – We aim at finding the value of an explanatory variable, through its expression in a large data vector, without knowing the link function between the explanatory variable and the data-space. Sliced Inverse Regression (SIR) method, allows the projection of a data-vector onto a subspace consistent with the explanatory variable variation. We suggest a method based on the SIR subspace, that gives the most efficient estimation of an unknown explanatory variable.

1 Introduction

L’évolution des détecteurs et de nos capacités de stockage de l’information font que la collecte des données astrophysiques est toujours plus massive. Des projets comme Gaia [1] ou encore le LSST [2] en sont le parfait exemple. Les bases de données ainsi constituées seront alors utilisées pour extraire le maximum d’informations quant à la caractérisation des objets observés. Ici, nous nous intéresserons à la question de la détermination des paramètres stellaires fondamentaux : température effective, gravité de surface et métallicité (ou composition chimique globale) d’une étoile à partir de la mesure de son spectre électromagnétique dans le visible. Cette détermination doit alors être faite sans connaissance d’un lien fonctionnel entre paramètres et données, mais en ayant à disposition soit une base de données de spectres déjà caractérisés, soit des spectres synthétiques. Il devient donc intéressant de mettre en œuvre des méthodes statistiques qui, à partir de telles bases de données, peuvent automatiquement déterminer les quantités physiques désirées sans que la relation entre ceux-ci et les données ne soit connue. SIR [3], avec une application encore très confidentielle dans un contexte astrophysique [4], permet d’obtenir une représentation des données cohérente avec le paramètre sans connaissance de la fonction qui les lie. Nous chercherons donc la manière la plus efficace d’utiliser les composantes du sous-espace (que nous appellerons "directions") obtenues grâce à SIR.

2 Sliced Inverse Regression

SIR recherche le sous-espace de représentation d’un vecteur de données x qui explique au mieux la valeur d’un para-

mètre y associé. Pour construire ce sous-espace à partir d’une base de données d’objets déjà caractérisés, on commencera par diviser l’espace des données en H tranches (sans recouvrement et contenant chacune le même nombre d’échantillons) suivant l’évolution du paramètre y auquel on s’intéresse. Ainsi les échantillons de la base de données sont regroupés en tranches dans lesquelles la valeur du paramètre y varie peu. Appliquer SIR permet de construire le sous-espace qui maximise la variance entre les tranches tout en gardant une variance globale réduite. De fait, on peut y voir un parallèle avec le fonctionnement de l’analyse factorielle discriminante [5] dans un contexte qui aurait été un contexte de classification.

2.1 Principe

Soit $\mathbf{X}$ la matrice des données dont chaque ligne est un échantillon associé à une valeur du paramètre y . Nous appellerons $\mathbf{X}_h$ la matrice des éléments de $\mathbf{X}$ qui appartiennent à la tranche h . Pour chaque tranche nous calculerons la moyenne $\bar{m}_h$ des lignes de $\mathbf{X}_h$.

La matrice $\mathbf{X}_H$ est la matrice des barycentres des tranches, dont chaque ligne est un des vecteur $\bar{m}_h$.

Γ la matrice de variance-covariance inter-tranches :

$$\Gamma = \frac{1}{H}(\mathbf{X}_H - \bar{\mathbf{X}}_H)^T(\mathbf{X}_H - \bar{\mathbf{X}}_H), \quad (1)$$

et Σ la matrice de variance-covariance totale :

$$\Sigma = \frac{1}{N}(\mathbf{X} - \bar{\mathbf{X}})^T(\mathbf{X} - \bar{\mathbf{X}}), \quad (2)$$

où N est le nombre d’échantillons de la base de donnée servant à construire le sous-espace, $\bar{\mathbf{X}}$ est la moyenne des lignes de $\mathbf{X}$ et

$\bar{\mathbf{X}}_H$ est la moyenne des lignes de $\mathbf{X}_H$. Les directions décrivant le sous-espace que l'on recherche sont les solutions de :

$$(\Sigma^{-1}\Gamma - \lambda \mathbf{I}_M)u = 0. \quad (3)$$

Les solutions de l'équation 3 correspondent aux vecteurs propres u associés aux plus grandes valeurs propres λ de $\Sigma^{-1}\Gamma$. M est la dimension de l'espace de départ.

Un point critique dans la mise en œuvre de SIR, est l'inversion de la matrice Σ . La structure des bases de données, fait que la matrice Σ n'est pas toujours de rang plein, elle est donc mal conditionnée. Lorsque c'est le cas, nous appliquerons, sur les données, une analyse en composante principale (PCA) [6] comme pré-traitement à SIR pour améliorer le conditionnement de Σ en diminuant la dimension de l'espace.

2.2 Exemple

Dans cet exemple l'espace de départ est de dimension 4, et nous chercherons à trouver le sous-espace de dimension 1 liant x à y sachant que $y(x) = 2x_1x_2 + x_3 + 0x_4 + \epsilon$ et que les éléments de x suivent une distribution uniforme $\mathcal{U}[0, 1]$. Le bruit $\epsilon \sim \mathcal{N}(0, 0.01I_4)$

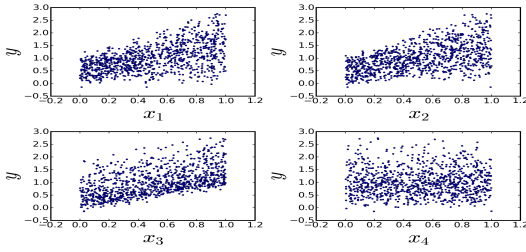


FIGURE 1 – Valeurs prises par y (ordonnées) par rapport aux réalisations de $\mathbf{X}$ (abscisses). Il n'est pas évident de trouver un lien entre les réalisations de $\mathbf{X}$ et la valeur prise par y à partir de cet espace.

La figure 1 montre, qu'en l'absence de la connaissance de la fonction liant x à y , il est nécessaire dans cet exemple de trouver une direction permettant de retrouver la valeur de y .

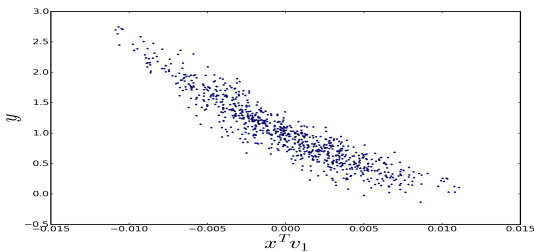


FIGURE 2 – Valeurs prises par y par rapport aux réalisations de $\mathbf{X}$ projetées sur la direction associée à la plus grande valeur propre donnée par SIR. Malgré la présence de bruit, on est en mesure d'observer un lien très clair entre les données et la valeur de y

Grâce à la figure 2, on peut voir que SIR, grâce à la prise en compte des variations de y lors de la recherche de directions pertinentes, est à même de fournir une projection qui montre une forte corrélation avec les données. L'application de SIR nous permet d'obtenir un sous-espace, formé par la ou les directions les plus pertinentes.

3 Détermination de paramètres grâce à SIR

Une fois que les données sont projetées il faut donc lier l'espace des données avec celui des paramètres. Le nombre optimal de directions doit donc être déterminé. Pour ce faire, nous pouvons procéder visuellement ou par un protocole de validation croisée, mais seulement dans l'hypothèse où les directions les plus pertinentes sont toujours les premières, et où la dimension optimale du sous-espace est toujours la même. Une fois tous les échantillons de la base projetés, l'échantillon dont on souhaite connaître la valeur du paramètre est à son tour projeté. L'estimateur de la valeur du paramètre est obtenue à partir de ces plus proches voisins dans le sous-espace, par une moyenne des valeurs du paramètre correspondant à ceux-ci. Dans cette partie nous essaierons de retrouver grâce à SIR les informations sur la profondeur et la largeur de raies spectrales gaussiennes bruitées :

$$x(\lambda; h, w) = 1 - h \times e^{-\frac{1}{2} \left(\frac{\lambda - \lambda_0}{w} \right)^2} + \epsilon \quad (4)$$

$\lambda_0 = 0$ est commune à tous les signaux. h et w sont uniformément distribués $h \sim \mathcal{U}[0.1, 0.9]$, $w \sim \mathcal{U}[1, 3]$ et le bruit ϵ suit une distribution gaussienne $\epsilon \sim \mathcal{N}(0, \sigma_\epsilon)$.

Dans ce cas, une PCA (sous-optimale) sera appliquée comme un pré-traitement à SIR de sorte à décroître suffisamment le conditionnement de Σ pour permettre son inversion. La variance du bruit supposée indépendante de la puissance du signal sur chaque composante de l'espace, est valable pour l'application qui nous concerne : des spectres haute résolution à haut rapport signal sur bruit.

3.1 Détermination du nombre optimal de voisins à considérer

Pour l'estimation de la valeur du paramètre, une moyenne au sens des k plus proches voisins a plus de sens qu'une régression linéaire, car pour ce qui concerne notre problème, le lien entre données et paramètres n'a pas de raison d'être linéaire. La considération des plus proches voisins permet de palier le problème des non-linéarités sous l'hypothèse que localement les valeurs du paramètre prises par les individus sont proches. Le choix du nombre de voisins dépend de la structure et de la richesse de la base de données de référence. Si l'on prend un trop grand nombre de voisins, on risque de ne plus pouvoir vérifier l'hypothèse de faible variabilité locale ; au contraire si l'on ne considère pas suffisamment de voisins, l'estimation sera très sensible au bruit. Dans le cadre d'une base de données échantillonnée de manière homogène dans les paramètres, il

est possible de déterminer un nombre optimal de voisins quelle que soit la position dans la base (pour avoir la meilleure précision possible compte tenu des non-linéarités). Pour évaluer ce nombre optimal de voisins, nous avons procédé par validation croisée.

3.2 Sélection des directions pertinentes

Nous émettons l'hypothèse que suivant la position de l'échantillon projeté dans le sous-espace les directions pertinentes pour sa représentation au sens du paramètre recherché ne sont pas nécessairement les mêmes. La variété portant l'information sur les variations du paramètre ne serait donc pas toujours portée par les mêmes directions du sous-espace. Nous proposons une alternative au sous-espace particulier en sélectionnant les directions jugées pertinentes pour chaque échantillon. Les directions choisies ne seront pas nécessairement celles des vecteurs propres associés aux plus grandes valeurs propres de $\Sigma^{-1}\mathbf{T}$, mais celles qui seront les plus informatives au sens de la parcimonie des valeurs prises par y compte tenu de la projection de l'échantillon à identifier (cf. [7]). Si l'on observe les directions obtenues lorsque l'on considère les largeurs, on obtient les résultats présentés figure 3.

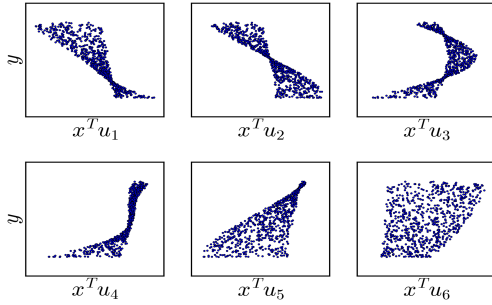


FIGURE 3 – Valeurs des paramètres de largeur des gaussiennes y par rapport aux valeurs des signaux $\mathbf{X}$ projetés sur les 6 premières directions données par SIR. On peut observer sur ces différents coefficients de projection, l'expression du lien non-linéaire entre les échantillons de la gaussienne et sa largeur.

La figure 3 montre que suivant la position de la projection d'un échantillon, les directions ne portent pas la même qualité d'information. Par exemple un échantillon se projetant sur les valeurs médianes de la direction u_1 obtient grâce à celle-ci une information très peu précise quant à la valeur de y . Au contraire si la valeur prise par $x^T u_1$ est proche de la limite supérieure, on est en mesure, grâce à cette seule direction, de déterminer précisément la valeur de y . On peut donc choisir suivant la position de l'échantillon dans le sous-espace quelles directions porteront l'information pertinente concernant la valeur de y associée. Dans le cas du calcul des profondeurs, le lien entre le paramètre et le signal étant linéaire, l'information pertinente est toujours portée par les mêmes directions de SIR. En procédant de la sorte, nous pouvons quantifier l'erreur moyenne en valeur absolue commise lors de la détermination des paramètres, et comparer ce résultat avec ce qu'aurait donné le meilleur sous-espace obtenu par PCA.

TABLE 1 – Erreurs moyennes en valeur absolue lors de la détermination des profondeurs et des largeurs des gaussiennes avec SIR ou la PCA.

	PCA	SIR
Profondeurs	0.0063	0.0025
Largeurs	0.029	0.024

Les erreurs présentées en table 1, sont calculées à partir de la détermination des valeurs de paramètres d'un cinquième de la base de données, le reste servant de base de référence et à la construction du sous-espace. Ces résultats présentés table 1 montrent que SIR permet d'obtenir une meilleure précision que ce que l'on peut attendre d'une PCA telle qu'utilisée dans [8].

4 Application à des spectres synthétiques

Nous comparerons les résultats obtenus avec ceux que l'on obtiendrait avec une PCA en suivant le protocole décrit dans [8]. Comme expliqué en section 2.1, une PCA sera appliquée en pré-traitement. La base de données contient 23000 spectres synthétiques et chacun de ces spectres est un vecteur de 8000 points couvrant le domaine spectral 390 – 680 nm (cf. [8]). Le pré-traitement ramènera cet espace à un espace de dimension 50 grâce à une PCA. En termes de valeurs prises par les paramètres (mêmes paramètres que dans [8]), la base est échantillonnée uniformément comme suit : $T_{\text{eff}} \in [4000 : 8000] K$, $\Delta(T_{\text{eff}}) = 100 K$, $\log(g) \in [4 : 5] \text{ dex}$, $\Delta(\log(g)) = 0.2 \text{ dex}$, $[Fe/H] \in [-1 : 0.5] \text{ dex}$, $\Delta([Fe/H]) = 0.1 \text{ dex}$, nous nous intéressons aussi à un quatrième paramètre, la vitesse de rotation projetée, $v \sin(i) \in [0 : 100] \text{ km/s}$, $\Delta(v \sin(i)) = 2 \text{ km/s}$ entre 0 km/s et 20 km/s et $\Delta(v \sin(i)) = 5 \text{ km/s}$ entre 20 km/s et 100 km/s. (Δ représente le pas d'échantillonnage de la base de données) Pour chaque calcul 1000 spectres sont aléatoirement extraits de la base et bruités avec un bruit gaussien centré et d'écart-type $\sigma = 0.03$ (cohérent avec les données observationnelles qui nous intéressent à terme) pour servir de base de test. On obtient alors avec SIR pour chacun des paramètres des directions telles que illustrés figure 4.

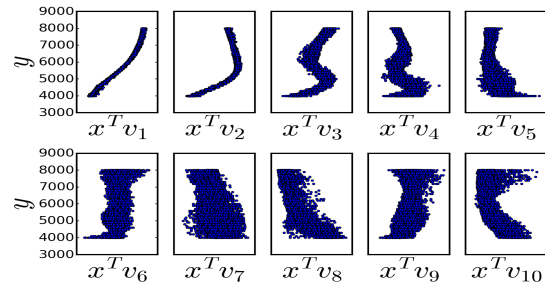


FIGURE 4 – Directions obtenues grâce à SIR pour la détermination des températures effectives (T_{eff})

On peut observer sur la figure 4 que les premières directions sont très informatives, de part le lien que l'on voit apparaître entre $x^T v_i$ et y , quant à la valeur prise par T_{eff} . Il est intéressant de noter que la direction numéro 5 par exemple n'apporte une

information pertinente que pour le tiers supérieur de $x^T v_5$.

On peut comparer les directions de la figure 4 avec les projections sur les composantes obtenues avec PCA en figure 5.

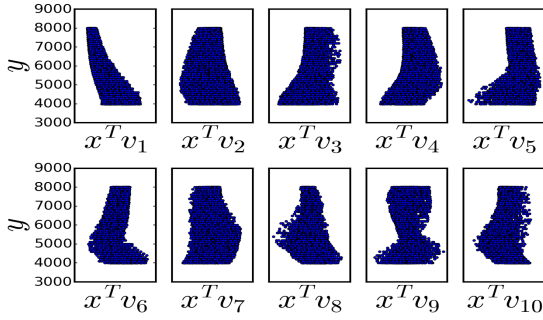


FIGURE 5 – Directions obtenues grâce à la PCA pour la détermination des températures effectives (T_{eff})

On s’aperçoit en regardant la figure 5, que les directions obtenues par PCA traduisent, de part leurs construction, beaucoup moins bien la variation du paramètre T_{eff} : cela signifie qu’il faut un plus grand nombre de directions pour avoir une information aussi précise qu’avec SIR. Mais prendre plus de directions signifie aussi plus de bruit. Il est donc peut probable que la PCA seule puisse avoir des résultats aussi précis que ceux atteints grâce à SIR.

Pour cet exemple, la base étant régulièrement échantillonnée pour tous les paramètres, on peut déterminer un nombre de voisins constant à considérer. Une dizaine de voisins permet de rester en dessous du nombre d’échantillons pour une valeur d’échantillonnage du paramètre, nous serons ainsi peu sensibles aux non linéarités. L’expérience montre que entre 10 et 30 voisins, les résultats ne sont pas différentiables

TABLE 2 – Erreurs moyennes en valeur absolue lors de la détermination des paramètres stellaires fondamentaux et du $v \sin(i)$ sur les spectres synthétiques.

	PCA	SIR
T_{eff}	132 K	61 K
$\log(g)$	0.25 dex	0.10 dex
$[Fe/H]$	0.07 dex	0.03 dex
$v \sin(i)$	1.2 km/s	0.09 km/s

On peut observer table 2 que, comme on s’y attendait, les résultats que l’on obtient avec SIR sont bien plus précis que ceux obtenus grâce à la PCA. L’inconvénient que l’on pourrait trouver à SIR, est que la méthode dans le cas où Σ est mal conditionnée, nécessite un pré-traitement. Elle nécessite aussi que l’on traite indépendamment chaque paramètre, ce qui rend son exécution plus lourde.

5 Conclusion

Nous avons vu que SIR permettait de projeter des données sur le sous-espace qui maximise la cohérence de variation entre

elles et un paramètre pré-défini. Cette approche ne nécessite pas la connaissance d’une fonction permettant de passer de l’espace des données à celui du paramètre. Il est nécessaire lorsqu’on l’applique de faire attention au conditionnement de Σ . Pour la détermination de paramètres stellaire, la méthode basée sur SIR se révèle être très efficace sur les spectres synthétiques par rapport à la méthode basée sur PCA. La PCA a l’avantage, quant à elle, de ne nécessiter qu’un espace de projection pour tous les paramètres. Les premiers tests sur des spectres observés montrent une amélioration d’environ 25% sur la température et sur la métallicité par rapport aux résultats de PCA. Cependant les bases de spectres observés étant souvent échantillonnées de manière non-homogène concernant certains paramètres -certaines plages de valeurs de $\log(g)$ contiennent très peu d’objets par exemple-, nous sommes confrontés à d’autres problèmes, notamment concernant la gravité de surface et la vitesse de rotation projetée. Un développement futur sera d’adapter la méthode pour des bases échantillonnées de manière non-homogène.

Références

- [1] The Gaia Collaboration. The gaia mission. *Astronomy & Astrophysics*, 595 :A1, 2016.
- [2] M. Juric and T. Tyson. Lsst data management : Entering the era of petascale optical astronomy. *Highlights of Astronomy*, 16 :675–676, March 2015.
- [3] K. C. Li. Sliced Inverse regression for dimension reduction. *Journal of the American Statistical Association*, 86 :316–327, 1991.
- [4] Caroline Bernard-Michel, Sylvain Douté, M Fauvel, Laurent Gardes, and Stephane Girard. Retrieval of Mars surface physical properties from OMEGA hyperspectral images using regularized sliced inverse regression. *Journal of geophysical research*, 114, 2009.
- [5] G. J. McLachlan. *Discriminant Analysis and Statistical Pattern Recognition*, ISSN 1940-6517 Wiley series in probability and mathematical statistics Volume 544 : Applied probability and statistics Wiley-Interscience paperback series. John Wiley & Sons, 2004.
- [6] I.T. Jolliffe. *Principal Component Analysis*. Springer Verlag, 1986.
- [7] V. Watson, J.-F. Trouilhet, F. Paletou, and S. Girard. Inference of an explanatory variable from observations in a high-dimensional space : Application to high-resolution spectra of stars. 2017 IEEE International workshop ECMSM, 2017.
- [8] F. Paletou, T. Böhm, V. Watson, and J.-F. Trouilhet. Inversion of stellar fundamental parameters from Espadons and Narval high-resolution spectra. *Astronomy & Astrophysics*, 573, 2015.

Chapitre 5

Étude comparative

Dans ce chapitre, il s'agira de comparer les différentes approches développées aux chapitres 3 et 4, autour de l'ACP et de SIR, avec les approches de références que sont les moindres carrés partiels (*Partial Least Squares*) PLS, au titre d'une méthode établie dans le domaine de la régression et MATISSE qui est une approche reconnue par le domaine de l'astrophysique pour le traitement des données de spectroscopie stellaire.

5.1 Méthodes de référence

Avant de se pencher sur la comparaison à proprement parler, nous présenterons dans cette partie les résultats de l'application des deux méthodes de référence, PLS et MATISSE, aux différents types de données.

5.1.1 Moindres carrés partiels

Pour commencer, nous allons appliquer la méthode PLS sur les données de l'exemple des raies gaussiennes dont nous souhaitons connaître les valeurs des dépressions et des largeurs (voir l'exemple du chapitre 3 paragraphe 3.3.3).

	RSB = 25 dB	RSB = 11 dB	RSB = 6 dB
Dépressions	0.0063	0.016	0.026
Largeurs	0.077	0.094	0.12

TABLE 5.1 – Résultats obtenus avec PLS pour l’estimation des dépressions et largeurs des gaussiennes pour différentes valeurs de RSB ($RSB = 10 \log_{10}(\frac{P_s}{P_b})$ où P_s est la puissance du signal et P_b la puissance du bruit).

Comme nous pouvions nous y attendre, PLS appliquant une régression linéaire globale, l’estimation des largeurs, paramètre non-linéaire, pose problème quel que soit le niveau de bruit (cf. table 5.1). La méthode ne pouvant pas prendre en compte les non-linéarités, ces dernières ont une telle influence sur l’erreur qu’il faudrait un niveau de bruit très élevé pour en voir l’effet. En revanche, l’approche par PLS permet d’avoir une assez bonne estimation des dépressions.

En appliquant PLS sur les données réelles on obtient les résultats suivants :

	T_{eff}	$\log(g)$	$[Fe/H]$	$v \sin(i)$
PLS standard	111	0.19	0.11	8.01
PLS local	109	0.16	0.088	6.19

TABLE 5.2 – Erreur d’estimation moyenne des trois paramètres fondamentaux T_{eff} (K), $\log(g)$ (dex), $[Fe/H]$ (dex) d’une part, et de la vitesse de rotation projetée $v \sin(i)$ (km/s) d’autre part. Application réalisée sur l’échantillon des 104 spectres du S4N considérés en utilisant la méthode PLS. La seconde ligne correspond aux résultats obtenus en appliquant les régressions locales plutôt que globales.

La table 5.2 montre que l’application de régressions locales permet de réduire (entre 2 et 25 % suivant le paramètre) l’erreur moyenne d’estimation pour tous les paramètres. En effet, ces paramètres se traduisent de manière non-linéaires dans les données et le principe des régressions locales est de palier aux non-linéarités.

5.1.2 MATISSE

Nous allons appliquer la méthode MATISSE, décrite dans Recio-Blanco et al. (2006), sur les mêmes jeux de données que les méthodes précédentes, tout d’abord sur les raies gaussiennes puis sur les données réelles.

Sur les données de l’exemple de gaussiennes bruitées, MATISSE a permis d’obtenir les résultats présentés table 5.3.

	RSB = 25 dB	RSB = 11 dB	RSB = 6 dB
Dépressions	0.0035	0.018	0.031
Largeurs	0.086	0.10	0.12

TABLE 5.3 – Résultats obtenus avec MATISSE pour l’estimation des dépressions et largeurs des gaussiennes pour différentes valeurs de RSB ($RSB = 10 \log_{10}(\frac{P_s}{P_b})$ où P_s est la puissance du signal et P_b la puissance du bruit).

La table 5.3 montre que MATISSE donne sur cet exemple des résultats du même ordre de grandeur que PLS. MATISSE donne tout de même une meilleure estimation des dépressions pour un faible niveau de bruit (colonne 1 ligne 1 du tableau). Cela confirme l’hypothèse selon laquelle MATISSE est surtout efficace pour l’estimation des paramètres linéaires, mais cela montre aussi la sensibilité au bruit à laquelle on peut s’attendre avec cette méthode.

En appliquant MATISSE aux données réelles, on obtient les résultats présentés table 5.4.

	T_{eff}	$\log(g)$	$[Fe/H]$	$v \sin(i)$
MATISSE	134	0.24	0.086	9.88

TABLE 5.4 – Erreur d’estimation moyenne des trois paramètres fondamentaux T_{eff} (K), $\log(g)$ (dex), $[Fe/H]$ (dex) et du $v \sin(i)$ (km/s), pour l’échantillon des 104 spectres du S4N considérés en utilisant la méthode MATISSE.

Les résultats de la table 5.4 montrent que MATISSE permet d'obtenir des résultats du même ordre de grandeur que PLS. Ces résultats sont plus globalement comparés aux résultats des autres méthodes dans la table 5.7.

5.2 Tables comparatives

On est à présent en mesure de dresser des tables comparatives concernant les différentes méthodes sur les données de test (gaussiennes dont l'on cherche les valeurs des largeurs et des dépressions paragraphe ; 3.3.3) et les données réelles grâce à la base Elodie et aux spectres S4N.

	ACP+k-PPV	SIR+k-PPV avec sélection	MATISSE	PLS local
RSB = 25db	0.0082	0.0043	0.0035	0.0042
RSB = 11 dB	0.014	0.013	0.018	0.014
RSB = 6 dB	0.026	0.025	0.031	0.023

TABLE 5.5 – Résultats obtenus avec les différentes méthodes pour l'estimation des dépressions des gaussiennes pour différentes valeurs de RSB.

La table 5.5 montre que les méthodes basées sur PLS et sur SIR donnent des résultats similaires, PLS devenant plus performante à 6 dB de RSB (MATISSE est par ailleurs plus performant à 25 dB). On peut en conclure que pour l'estimation des paramètres linéaires les deux méthodes se valent. La méthode basée sur l'ACP est elle aussi compétitive sauf à haut rapport signal sur bruit.

	ACP+k-PPV	SIR+k-PPV avec sélection	MATISSE	PLS local
RSB = 25db	0.020	0.016	0.086	0.037
RSB = 11 dB	0.049	0.042	0.10	0.065
RSB = 6 dB	0.080	0.078	0.12	0.077

TABLE 5.6 – Résultats obtenus avec les différentes méthodes pour l’estimation des largeurs des gaussiennes pour différentes valeurs de RSB ($RSB = 10 \log_{10}(\frac{P_s}{P_b})$ où P_s est la puissance du signal et P_b la puissance du bruit).

La table 5.6 montre que l’approche basée sur SIR donne les meilleurs résultats (sauf à faible RSB où l’ACP, SIR et PLS sont équivalents). On peut déduire que l’approche basée sur SIR avec la sélection des directions pertinentes est un bon compromis pour traiter des paramètres linéaires et non-linéaires, pour des niveaux de bruits variables. Là où MATISSE est surtout efficace pour un grand RSB, PLS est surtout efficace quand le paramètre traité est linéaire, et l’ACP est surtout efficace lorsque la puissance du bruit est importante.

	ACP+k-PPV	SIR+k-PPV avec sélection	MATISSE	PLS local
T_{eff} (K)	102	84	134	109
$\log(g)$ dex	0.20	0.17	0.24	0.16
$[Fe/H]$ dex	0.087	0.068	0.086	0.088
$v \sin(i)$ km/s	1.37	1.38	9.88	6.19

TABLE 5.7 – Résultats obtenus avec les différentes méthodes pour l’estimation des paramètres fondamentaux et du $v \sin(i)$ des spectres du S4N à partir des spectres Elodie.

La table 5.7 montre que l’approche basée sur SIR, donne systématiquement les erreurs les plus faibles. Pour la vitesse de rotation projetée, $v \sin(i)$, on obtient le même résultat qu’avec l’ACP, et pour la gravité de surface, $\log(g)$, on obtient le même résultat (à 5%

près) qu’avec PLS. Les paramètres sur lesquels on peut voir une nette amélioration sont la température effective et la métallicité.

Au chapitre 4, nous avons évoqué que le problème pour l’estimation de la vitesse de rotation projetée était lié à l’échantillonnage de la base de référence. L’estimation de la vitesse de rotation projetée ne pose pas de problème dans le cas des spectres synthétiques, car ceux-ci sont échantillonnés de manière dense. Dans le cas des spectres réels, seuls quelques spectres ont une vitesse de rotation projetée supérieure à 25 km/s. La sous-représentativité des spectres à haute vitesse de rotation projetée dans les données réelles explique la différence avec les résultats sur les spectres synthétiques. Les résultats concernant l’estimation de la gravité de surface ne sont pas aussi bons que ce que l’on peut trouver dans la littérature¹, mais les mesures spectroscopiques ne représentent pas le meilleur outil pour une estimation précise de ce paramètre².

Pour ce qui concerne la température, l’annexe 1 montre que notre estimation est très proche de la valeur de référence par rapport aux valeurs issues de la littérature.

Nous pouvons conclure de cette étude comparative que la méthode proposée, qui consiste en **l’application de SIR avec une sélection des directions pertinentes et une méthode des k-PPV, donne quasi-systématiquement les meilleurs résultats sur les données astrophysiques**. La sélection des directions n’apporte une amélioration que sur des données à haut rapport signal sur bruit, mais n’engendre pas de détérioration lorsque le bruit augmente.

Cette nouvelle approche représente donc, pour l’estimation des paramètres stellaires fondamentaux, une alternative intéressante aux méthodes déjà utilisées par la communauté de l’astrophysique. Cette alternative a déjà son exploitation concrète dans le domaine de l’astrophysique grâce à Kassounian et al. (2018) présenté en annexe.

1. cf. annexe 1.

2. L’astéro-sismologie permet des mesures plus fiables et plus précises de $\log(g)$.

Conclusion

Dans ce mémoire, nous avons analysé plusieurs approches de traitement statistique de données qui ont été appliquées à des problèmes simples puis à des données de spectroscopie stellaire relatives à l'application ayant motivé ce travail de thèse.

Au chapitre 3, partant de l'application d'une réduction de dimensionnalité à l'aide d'une ACP couplée à une méthode des k -plus proches voisins, nous avons proposé des solutions pour améliorer les performances de la méthode, compte tenu de la nature non-linéaire des données que nous cherchons à traiter. C'est ainsi que nous avons proposé une méthode visant à définir une localité pour approcher les non-linéarités avec des régressions linéaires. Nous avons aussi proposé une approche plus novatrice basée sur la sélection de directions post-projection afin, notamment, de s'extraire de l'aspect non-linéaire du lien entre l'espace de projection et l'espace du paramètre que l'on cherche à identifier. Cette sélection des directions repose sur l'hypothèse que le sous espace optimal ne nécessite pas l'ensemble des directions pour tous les individus. Cette hypothèse a ensuite été validée sur des exemples avant que la méthode de sélection des directions ne soit appliquée sur des données astrophysiques.

Ces approches n'ont montré que peu d'impact sur les résultats obtenus par la régression basée sur l'ACP. En revanche, au chapitre 4 nous avons montré que l'utilisation de SIR, en plus de donner de meilleurs résultats, est plus sensible à la sélection des directions qui lui permet, dans le cas d'un haut rapport signal sur bruit, une réduction conséquente de l'erreur d'estimation. L'emploi de SIR permet ainsi de réduire globalement l'erreur d'estimation, et la sélection des directions permet encore d'améliorer ces résultats.

Du chapitre 5, nous pouvons conclure que la méthode de régression basée sur SIR que nous proposons est pertinente du point de vue de l'application. En effet, celle-ci permet d'obtenir les plus faibles erreurs d'estimation. Dans les quelques cas pour lesquels elle ne donne pas les meilleurs résultats, elle donne des erreurs très légèrement plus élevées que les plus faibles obtenues par ailleurs. Ainsi, sur les données réelles, l'approche proposée permet une réduction de l'erreur moyenne d'estimation d'environ 20% pour l'estimation des températures effectives et des métallicités. Dans le cas de la gravité de surface, la méthode basée sur les moindres carrés partiels locaux, permet d'avoir une erreur moyenne 5% inférieure à l'alternative que nous proposons (cette dernière permettant une réduction de 25% de l'erreur par rapport aux méthodes précédemment en usage dans le domaine d'application). Pour ce qui concerne la vitesse de rotation projetée, les résultats sont sensiblement les mêmes que pour la régression basée sur l'analyse en composantes principales de (Paletou et al., 2015a).

Ces résultats sont perfectibles, mais un enjeu majeur pour réduire les erreurs d'estimation sera lié à la structure de la base de données de référence.

Une perspective de travail futur pourra se concentrer sur l'optimisation d'une base de données de spectres synthétiques dont la densité en termes d'échantillonnage des paramètres est optimale. Du point de vue des développements liés aux méthodes, la sélection des directions permet, avec une méthode de projection linéaire, d'optimiser le sous-espace de projection pour chaque spectre. Cela permet parfois de lever des non-linéarités, comme cela a été montré au chapitre 3, mais cela permet aussi de réduire l'impact du bruit lorsque, pour un individu donné, l'information n'est pas portée par toutes les directions du sous-espace. La sélection des directions se base sur la variance empirique comme indicateur de la densité de probabilité du paramètre, en fonction de chacune des directions.

Une autre perspective pourrait consister en la détermination précise de densité de probabilités conditionnelles du paramètre en fonction des différentes directions à partir d'histogrammes, même si le calcul de ces histogrammes nécessitera une base de données de spectres plus dense que celles utilisées.

Bibliographie

- Allende Prieto, C., Barklem, P. S., Lambert, D. L., & Cunha, K. (2004). S⁴N : A spectroscopic survey of stars in the solar neighborhood. The Nearest 15 pc. *Astronomy and Astrophysics*, 420, 183–205.
- Babenko, A., Slesarev, A., Chigorin, A., & Lempitsky, V. (2014). *Neural Codes for Image Retrieval*. (p. 584-599), Springer International Publishing.
- Bailer-Jones, C. A. L. (2000). Stellar parameters from very low resolution spectra and medium band filters. T_{eff} , $\log(g)$ and $[M/H]$ using neural networks. *Astronomy and Astrophysics*, 357, 197-205.
- Bailer-Jones, C. A. L. (2013). The Gaia astrophysical parameters inference system (Apsis). *Astronomy and Astrophysics*, 559, A74.
- Bailer-Jones, C. A. L., Irwin, M., & Von Hippel, T. (1998). Automated classification of stellar spectra — II. two-dimensional classification with neural networks and principal components analysis. *Monthly Notices of the Royal Astronomical Society*, 298(2), 361–377.
- Bandžuch, P., Morhác, M., & Krištiak, J. (1997). Study of the Van Cittert and Gold iterative methods of deconvolution and their application in the deconvolution of experimental spectra of positron annihilation. *Nuclear Instruments and Methods in Physics Research Section A : Accelerators, Spectrometers, Detectors and Associated Equipment*, 384(2), 506 – 515.

- Baranne, A., Queloz, D., Mayor, M., Adrianzyk, G., Knispel, G., Kohler, D., Lacroix, D., Meunier, J.-P., Rimbaud, G., & Vin, A. (1996). ELODIE : a spectrograph for accurate radial velocity measurements. *Astron. Astrophys. Suppl. Ser.*, 119(2), 373–390.
- David Dussault, P. H. (2004). Noise performance comparison of ICCD with CCD and EMCCD cameras. *Proc.SPIE*, 5563, 10.
- De La Hoz, E., De La Hoz, E., Ortiz, A., Ortéga, J., & Prieto, B. (2015). PCA filtering and probabilistic SOM for network intrusion detection. *Neurocomputing*, 164(10), 71–81.
- Deeming, T. J. (1964). Stellar spectral classification : I. application of component analysis. *Monthly Notices of the Royal Astronomical Society*, 127(6), 493–516.
- Donati, G. B. (1863). On the Striæ of Stellar Spectra, from the *Memorie Astronomiche*. *Monthly Notices of the Royal Astronomical Society*, 23, 100.
- Freitag, D. (2000). Machine learning for information extraction in informal domains. *Machine Learning*, 39(2), 169–202.
- Gebran, M., Farah, W., Paletou, F., Monier, R., & Watson, V. (2016). A new method for the inversion of atmospheric parameters of A/Am stars. *Astronomy and Astrophysics*, 589, A83.
- Goebel, J., Stutz, J., Volk, K., Walker, H., Gerbault, F., Self, M., Taylor, W., & Cheeseman, P. (1989). A bayesian classification of the IRAS LRS Atlas. *Astronomy and Astrophysics*, 222, no. 1-2, L5-L8.
- Gulati, R. K., Gupta, R., Gothoskar, P., & Khobragade, S. (1994). Stellar spectral classification using automated schemes. *Astrophysical Journal, Part 1*, 426, no. 1, 340-344.
- Gustafsson, B., Bell, R. A., Eriksson, K., & Nordlund, A. (1975). A grid of model atmospheres for metal-deficient giant stars. I. *Astronomy and Astrophysics*, vol. 42, no. 3, 407-432.
- Hastie, T., & Stuetzle, W. (1989). Principal curves. *J. Am. Stat. Assoc.*, 84(406), 502–516.

- Henning, J. L. (2000). SPEC CPU2000 : measuring CPU performance in the new millennium. *Computer*, 33(7), 28–35.
- Hubeny, I., & Lanz, T. (1992). Accelerated complete-linearization method for calculating NLTE model stellar atmospheres. *Astronomy and Astrophysics*, 262, 501–514.
- Jolliffe, I. (1986). *Principal Component Analysis*. Springer Verlag.
- Kaipio, J. P., & Somersalo, E. (2005). *Statistical and Computational Inverse Problems, Chapter 2 : Classical Regularization Methods*. New York, NY : (p. 7-48), Springer New York.
- Kassounian, S., Gebran, M., Paletou, F., & Watson, V. (2018). Sliced Inverse Regression : application to stellar fundamental parameters. *Monthly Notices of the Royal Astronomical Society (submitted)*.
- Kataria, A., & Singh, M. (2013). A Review of Data Classification Using K-Nearest Neighbour Algorithm. *International Journal of Emerging Technology and Advanced Engineering*, 3, 354.
- Katz, D., Soubiran, C., Cayrel, R., Adda, M., & Cautain, R. (1998). On-line determination of stellar atmospheric parameters T_{eff} , $\log(g)$, $[\text{Fe}/\text{H}]$ from ELODIE echelle spectra. I. the method. *Astronomy and Astrophysics*, 338, 151-160.
- Kurucz (1970). Model atmosphere codes : ATLAS12 and ATLAS9. *SAO special report 309, Cambridge, USA*.
- Kurucz, R. L. (2005). ATLAS12, SYNTHE, ATLAS9, WIDTH9, et cetera. *Memorie della Societa Astronomica Italiana Supplementi*, 8, 14.
- Li, K. C. (1991). Sliced Inverse regression for dimension reduction. *Journal of the american staistical association*, 86, 316–327.
- McLachlan, G. J. (2004). *Discriminant Analysis and Statistical Pattern Recognition Volume 544 of Wiley Series in Probability and Statistics, Wiley series in probability and*

mathematical statistics : Applied probability and statistics Wiley-Interscience paperback series. John Wiley & Sons.

Morgan, W. W., Keenan, P. C., & Kellman, E. (1943). *An atlas of stellar spectra, with an outline of spectral classification.* Chicago, Ill., The University of Chicago press [1943].

Orlov, S. S., Phillips, W., Bjornson, E., Takashima, Y., Sundaram, P., Hesselink, L., Okas, R., Kwan, D., & Snyder, R. (2004). High-transfer-rate high-capacity holographic disk data-storage system. *Appl. Opt.*, 43(25), 4902–4914.

Paletou, F., Böhm, T., Watson, V., & Trouilhet, J.-F. (2015a). Inversion of stellar fundamental parameters from Espadons and Narval high-resolution spectra. *Astronomy & Astrophysics*, 573, A67.

Paletou, F., Gebran, M., Houdebine, E. R., & Watson, V. (2015b). Principal component analysis-based inversion of effective temperatures for late-type stars. *Astronomy and Astrophysics*, 580, A78.

Pennycook, S., Varela, M., Hetherington, C., & Kirkland, A. (2006). Materials advances through aberration-corrected electron microscopy. *MRS Bulletin*, 31(1), 36–43.

Petit, D., & Maillet, D. (2008). Techniques inverses et estimation de paramètres. partie 2. *Techniques de l'ingénieur Physique statistique et mathématique*, 42619210.

Recio-Blanco, A., Bijaoui, A., & de Laverny, P. (2006). Automated derivation of stellar atmospheric parameters and chemical abundances : the MATISSE algorithm. *Monthly Notices of the Royal Astronomical Society*, 370, 141-150.

S. Perruchot, e. a. (2008). The SOPHIE spectrograph : design and technical key-points for high throughput and high stability. *Proc. SPIE*, 7014, 12.

Saha, N. M. (1921). On a physical theory of stellar spectra. *Proceedings of the Royal Society of London A : Mathematical, Physical and Engineering Sciences*, 99(697), 135–153.

- Saporta, G. (2011). *Probabilités, analyse des données et statistique, Chapitre 7*. Editions Technip.
- Tenenhaus, M. (1998). *La régression PLS, théorie et pratique, Chapitre 7*. Editions Technip.
- Thévenin, F., Bijaoui, A., & Katz, D. (2003). Chemical abundances from GAIA spectra. In U. Munari (Ed.) *GAIA Spectroscopy : Science and Technology*, vol. 298 of *Astronomical Society of the Pacific Conference Series*, (p. 291).
- von Hippel, T., Storrie-Lombardi, L. J., Storrie-Lombardi, M. C., & Irwin, M. J. (1994). Automated classification of stellar spectra – I. initial results with artificial neural networks. *Monthly Notices of the Royal Astronomical Society*, 269(1), 97–104.
- Watson, V., Trouilhet, J.-F., Paletou, F., & Gebran, M. (2017b). *Sliced Inverse Regression* pour l’estimation de paramètres stellaires fondamentaux. Grets - 2017.
- Watson, V., Trouilhet, J.-F., Paletou, F., & Girard, S. (2017a). *Inference of an explanatory variable from observations in a high-dimensional space : Application to high-resolution spectra of stars*. 2017 IEEE International Workshop of Electronics, Control, Measurement, Signals and their application to Mechatronics (ECMSM).
- Whitney, C. A. (1983). Stellar acoustics. I - adiabatic pulse propagation and modal resonance in polytropic models of bump cepheids. *Astrophysical Journal*, 274, 830-839.
- Wold, S., Martens, H., & Wold, H. (1983). *The multivariate calibration problem in chemistry solved by the PLS method*. (p. 286-293) Berlin, Heidelberg : Springer Berlin Heidelberg.

Annexes

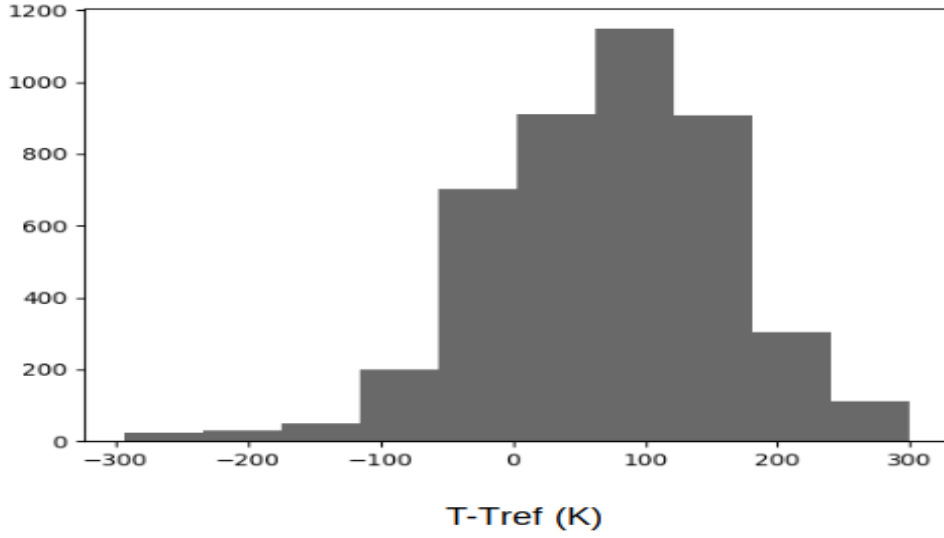


FIGURE 5.1 – Histogramme de la répartition des valeurs de températures effectives des catalogues VizieR autour des valeurs de référence pour les spectres du S4N utilisés. En abscisse est représenté la valeur de température, en ordonnées le nombre de valeurs par intervalle. $\sigma = 89\ K$.

Annexe 1 : Évaluation de l’incertitude des labels du S4N

La base de données de test, dans le cas des données réelles utilisées dans ce travail, recense des objets dont les paramètres fondamentaux (température effective, gravité de surface et métallicité) ont aussi été déterminés par divers auteurs, parfois nombreux suivant les étoiles. De cela, une valeur pour chaque paramètre fondamental, a été retenue pour chaque étoile. Pour les valeurs de références du S4N, nous utilisons celles fournies par Allende Prieto et al. (2004), mais on peut étudier la fiabilité de cette valeur en la comparant à toutes celles issues de la bibliographie disponible pour le même objet³. Ainsi nous pourrions, en définitive, avoir une idée de la précision typiquement atteinte quelque soit la méthode de traitement utilisée.

Nous représentons dans la figure 5.1 la dispersion moyenne des valeurs de la température

3. Catalogues VizieR : <http://cdsweb.u-strasbg.fr/>

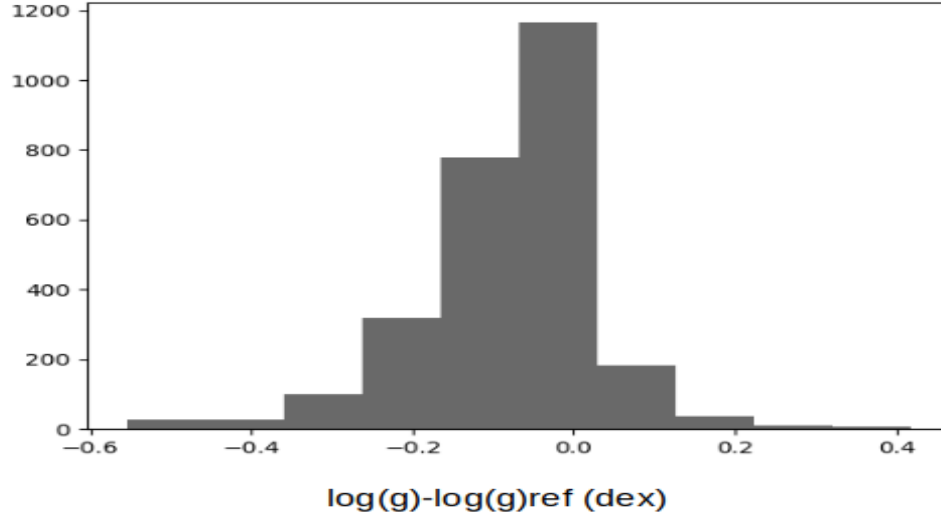


FIGURE 5.2 – Histogramme de la répartition des valeurs de gravité de surface autour des valeurs de référence pour les spectres du S4N utilisés. En abscisse est représenté la valeur de gravité, en ordonnées le nombre de valeurs par intervalle. $\sigma = 0.12 \text{ dex}$.

effective T_{eff} autour de la valeur retenue donnée par Allende Prieto et al. (2004). On observe une répartition qui a un écart type $\sigma = 89 \text{ K}$, cela signifie que la plupart (66%) des méthodes s'accordent sur la valeur de température à 89 K près.

La figure 5.2 représente la répartition des valeurs de gravité de surface. Ici l'écart type de la répartition est de $\sigma = 0.12 \text{ dex}$.

La figure 5.3 représente la répartition des valeurs de métallicité. Ici l'écart type de la répartition est de $\sigma = 0.24 \text{ dex}$.

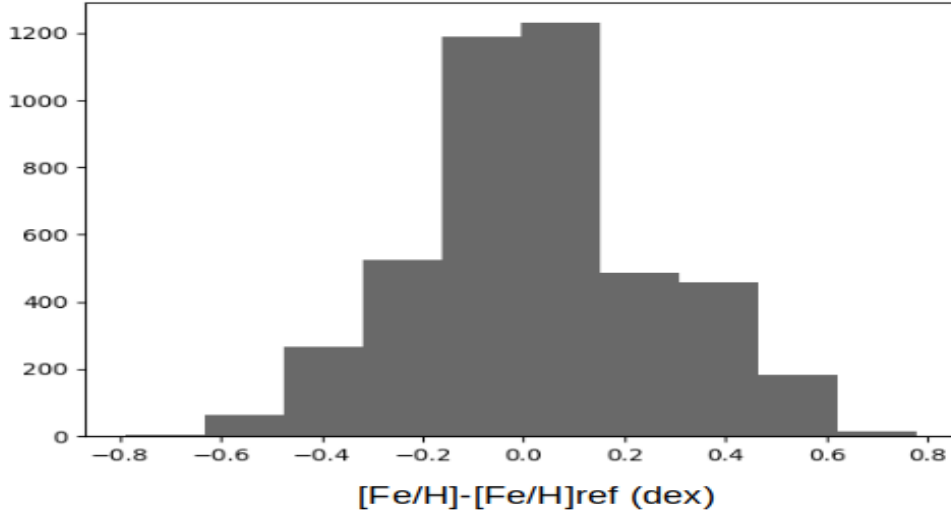


FIGURE 5.3 – Histogramme de la répartition des valeurs de métallicité autour des valeurs de référence pour les spectres du S4N utilisés. En abscisse est représenté la valeur de métallicité, en ordonnées le nombre de valeurs par intervalle. $\sigma = 0.24 \text{ dex}$.

Les résultats présentés ci-dessus montrent la disparité des valeurs obtenues, dans la littérature, pour l'estimation des paramètres stellaires fondamentaux.

Les valeurs de σ présentées permettent de quantifier cette disparité. On montre alors que 66% des méthodes s'accordent sur des valeurs de températures effectives à 89 K près (entre $-\sigma$ et $+\sigma$). Pour les valeurs de métallicité, ces valeurs s'accordent à 0.24 dex près et pour des valeurs de gravité de surface à 0.12 dex près. La dispersion concernant la métallicité est surprenante étant donné que les méthodes que nous avons utilisées dans cette étude nous permettent d'être beaucoup plus près de la valeur donnée par Allende Prieto et al. (2004).

Cette annexe montre que la méthode que nous proposons permet l'obtention de résultats cohérents avec la distribution statistique des valeurs de la littérature, sauf dans le cas de la gravité de surface. Ce manque de précision concernant la gravité de surface vient probablement de l'utilisation exclusive de données spectroscopiques de notre part.

Les données représentées ici agrègent d'autres types de données (photométriques ou astéro-interférométriques), ce qui explique vraisemblablement le taux de précision de 0.12 dex.

Annexe 2 : Autres contributions

Principal component analysis-based inversion of effective temperature for late-type stars (Research Note)

Résumé : Nous montrons que le domaine d'application de la méthode d'inversion de (Paletou et al., 2015a) basée sur l'analyse en composantes principales peut être étendue à des données d'étoiles froides. Au delà du fait que cela représente une expansion du domaine original d'application, qui concernait les étoiles de type FGK, nous avons aussi utilisé des spectres synthétiques pour la base de données d'apprentissage. Nous discutons nos résultats concernant la température effective, avec les évaluations préalables rendues accessibles grâce aux services VizieR et Simbad du CDS.

Principal component analysis-based inversion of effective temperatures for late-type stars[★] (Research Note)

F. Paletou^{1,2}, M. Gebran³, E. R. Houdebine^{1,2,4}, and V. Watson^{1,2}

¹ Université de Toulouse, UPS-Observatoire Midi-Pyrénées, Irap, 31028 Toulouse Cedex 4, France
 e-mail: fpaletou@irap.omp.eu

² CNRS, Institut de Recherche en Astrophysique et Planétologie, 14 Av. E. Belin, 31400 Toulouse, France

³ Department of Physics and Astronomy, Notre Dame University-Louaize, PO Box 72, Zouk Mikael, Lebanon

⁴ Armagh Observatory, College Hill, BT61 9DG Armagh, Northern Ireland

Received 24 June 2015 / Accepted 14 July 2015

ABSTRACT

We show how the range of application of the principal component analysis-based inversion method of Paletou et al. (2015, A&A, 573, A67) can be extended to data for late-type stars. Besides being an extension of its original application domain, which applied to FGK stars, we also used synthetic spectra for our learning database. We discuss our results for effective temperatures in comparison with previous evaluations made available from VizieR and Simbad services at CDS.

Key words. virtual observatory tools – stars: fundamental parameters – stars: late-type

1. Introduction

Effective temperatures of late-type stars were inverted from HARPS spectra, using the principal component analysis-based (PCA) method detailed in Paletou et al. (2015). In the latter study, fundamental parameters of FGK stars were inverted using a so-called learning database generated from the Elodie stellar spectra library (see Prugniel et al. 2007) using observed spectra for which fundamental parameters were already evaluated. Also, spectra considered in the Paletou et al. (2015) study had typical spectral resolution $\mathcal{R}$ of 50 000 (Allende Prieto et al. 2004) and 65 000 (Petit et al. 2014), i.e. values significantly lower than HARPS data.

In this study, the inversion of the effective temperature, T_{eff} , from spectra of late-type (dwarf) stars of K and M spectral types is performed using a database of synthetic spectra. We discuss hereafter comparisons with published values collected from the VizieR and Simbad services of CDS.

2. The learning database

A grid of 6336 spectra was computed using S₄₈ synthetic spectra code (Hubeny & Lanz 1992) and Kurucz A₁₂ model atmospheres (Kurucz 2005). The linelist was built from Kurucz (1992) `gfhyperl1.dat`¹.

For our purposes, we adopted a grid of parameters such that T_{eff} is in a 3500–4600 K range with a 100 K step, $\log g$ is in the range of 4–5 dex with a 0.2 dex step, metallicity [Fe/H] is in a -2 – $+0.5$ range with a 0.25 dex step and, finally, $v \sin i$ varies from 0 to 14 km s^{−1} with a 2 km s^{−1} step.

For all models, the microturbulent velocity was fixed at $\xi_t = 1$ km s^{−1} and $[\alpha/\text{Fe}]$ was set to 0. The spectral resolution of the HARPS spectrograph, i.e. $\mathcal{R} = 115\,000$ was adopted for the production of this set of synthetic spectra. We finally limited the study to a spectral band centred around the Na D-doublet, ranging from 585.3 to 593.2 nm.

Even though the choice of Atlas for such cool stars as well as the choice of the spectral bands we considered may be questionable, we show that effective temperatures we inverted from HARPS data are realistic enough for several further studies (e.g. rotation-activity correlations vs. the spectral type).

3. Observational and reference data

We used 57 high-resolution HARPS spectra taken from the ESO Science Archive Facility. The selection targeted late-type dwarf K stars, and early-type dwarf M stars. We also gave preference to high signal-to-noise spectra. Related objects are listed in Fig. 1; see also Table 1.

All reference data were collected from VizieR catalogues, using the A² Python modules already mentioned by Paletou & Zolotukhin (2014).

4. Results

Given the “bulk” of nearest neighbours we consider for our inversion procedure, we estimate that 150 K is a typical upper value for the uncertainty in our derived effective temperature.

[★] Based on data obtained from the ESO Science Archive Facility.

¹ <http://kurucz.harvard.edu>

² astroquery.readthedocs.org

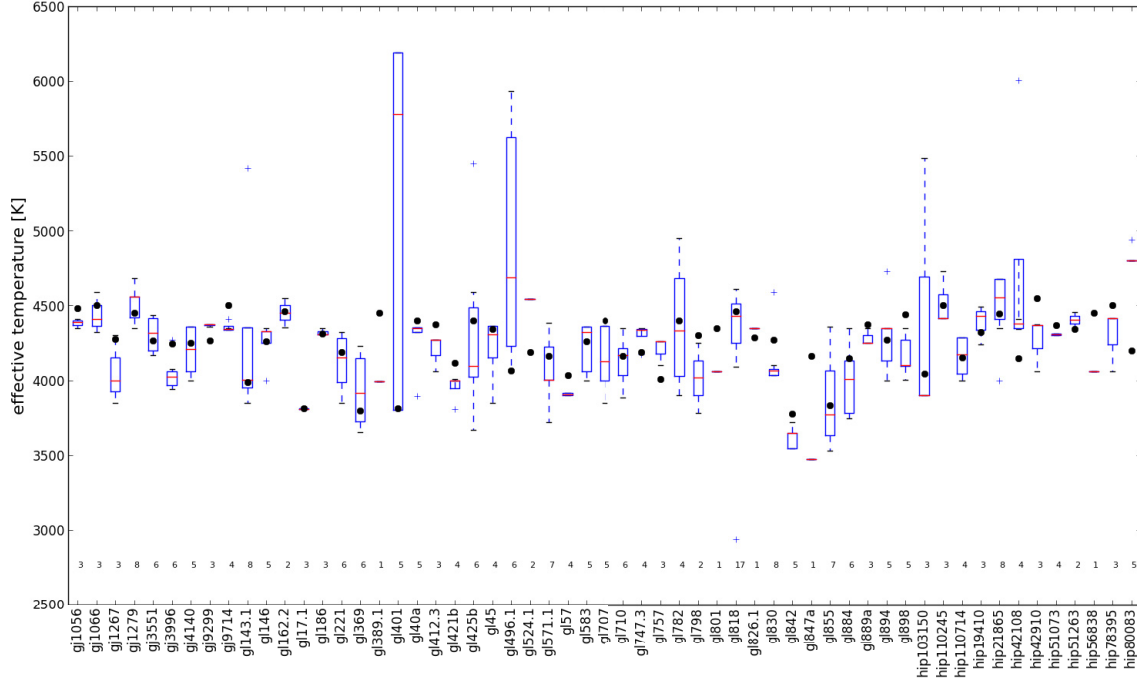


Fig. 1. Comparison between our estimate of effective temperatures ($\bullet$), and the values we got from available VizieR catalogues. The latter collections are represented as classical boxplots. (The horizontal bar inside each box indicates the median, or Q_2 value, while each box extends from first quartile, Q_1 , to third quartile Q_3 . Extreme values, still within a 1.5 times the interquartile range away from either Q_1 or Q_3 , are connected to the box with dashed lines. Outliers are denoted by a “+” symbol.) Objects we studied are listed along the horizontal axis. In addition, for each object at level $T_{\text{eff}} \sim 2800$ K, we explicated the number of values found among all VizieR catalogues.

From Fig. 1, we can identify eight major outliers considering deviations from our estimate of the effective temperature and values published and made available through the VizieR service at CDS.

GL 401 is a M3V star. It is also an interesting case, from the point of view of fundamental parameters made available. We retrieved five determinations of its effective temperature, spanning an impressive range of values, as can be seen from the corresponding boxplot in Fig 1. However, our estimate of 3810 K is in excellent agreement with the recent value of 3804 K derived by Gaidos et al. (2014).

The case of the K9V star GL496.1 is very similar. We extracted six values from VizieR, and again, the most recent value of 4075 K given by Gaidos et al. (2014) is in excellent agreement with our estimate of 4063 K. We also note the very different value of 4685 K given by Santos et al. (2013), and reported at Simbad at CDS for this object.

GL 389.1 is believed to be a K5.5V star (Gray et al. 2006). We could only identify only one data point at VizieR, an estimate of 3990 K given by Lafrasse et al. (2010, also catalogue id. II/320). Our own estimate is significantly hotter at 4450 K, which appears to be more consistent with the spectral type found at Simbad.

According to Simbad at CDS, GL 847 A is a K4 star (Van Leeuwen 2007). This is clearly not consistent with our estimate of an effective temperature of 4160 K, nor with the only value we could retrieve from VizieR, i.e. the estimate by Morales et al. (2008) yielding 3470 K. As for the precedent object, the lack of data available in catalogues makes the assessment of a reliable reference value difficult.

HIP 42108 is a K6V star (Gray et al. 2006) for which we could estimate an effective temperature of 4147 K. From

available catalogues, we found the most recent estimate of 4343 K given by McDonald et al. (2012), which is also very close to the alternative estimate of Wright et al. (2003). Lafrasse et al. (2010) also report a very close estimate of 4410 K. For this object, however, Ammons et al. (2006) report a very different value of 6005 K.

A similar case of overestimation has been identified for GL 524.1. We obtained an effective temperature of 4186 K, while Ammons et al. (2006) provide a significantly larger value of 4554 K. Unfortunately, we could not find alternative estimates for this parameter for this object (there is an entry for GL 524.1 reported by Morales et al. 2008, but no T_{eff} value is given).

HIP 56838 is, according to Simbad, a K6V star based on the classification of Gray et al. (2006). Our estimate of its effective temperature is 4200 K, in agreement with a K6 (main sequence) spectral type, while the only VizieR data we could get is 4060 K (Wright et al. 2003). Even though this object can also be found in the Ammons et al. (2006) catalogue, no T_{eff} value is given there. Besides, a much cooler temperature of 3800 K, which would better correspond to a M0V spectral type, was recently given by Kordopatis et al. (2013).

HIP 80083 is quite consistently given at 4800 K (Sousa et al. 2011; Adibekyan et al. 2012; Carretta 2013), or even hotter (Ammons et al. 2006). Our inversion procedure gives an effective temperature significantly lower at 4200 K, typical of a spectral type later than K4, as indicated by Simbad at CDS.

Table 1 summarizes our results. It displays our inverted $T_{\text{eff}}^{(\text{inv.})}$ and two reference values. The first one, $T_{\text{eff}}^{(\text{clos.})}$, was defined as the value found in VizieR catalogues closest to our estimate, while $T_{\text{eff}}^{(\text{med.})}$ is the median of all catalogue values.

Table 1. Inverted and reference effective temperatures for all objects.

Object	$T_{\text{eff}}^{(\text{inv.})}$ [K]	$T_{\text{eff}}^{(\text{clos.})}$ [K]	$T_{\text{eff}}^{(\text{med.})}$ [K]
gj1056	4480.0	4410.0	4391.0
gj1066	4500.0	4590.0	4410.0
gj1267	4275.0	4301.0	4000.0
gj1279	4450.0	4424.0	4556.0
gj3551	4267.0	4287.0	4318.5
gj3996	4244.0	4272.0	4023.0
gj4140	4250.0	4210.0	4210.0
gj9299	4267.0	4360.0	4372.0
gj9714	4500.0	4410.0	4344.5
gl143.1	3985.0	3970.0	4001.0
gl146	4260.0	4250.0	4329.0
gl162.2	4462.5	4550.0	4452.5
gl17.1	3812.5	3809.0	3809.0
gl186	4311.0	4307.0	4307.0
gl221	4186.0	4282.0	4151.0
gl369	3796.0	3915.0	3915.0
gl389.1	4450.0	3990.0	3990.0
gl401	3810.0	3804.0	5780.0
gl40a	4400.0	4352.0	4350.0
gl412.3	4375.0	4271.0	4271.0
gl421b	4117.0	4008.0	3995.0
gl425b	4400.0	4590.0	4097.5
gl45	4340.0	4361.0	4305.5
gl496.1	4063.0	4075.0	4685.0
gl524.1	4186.0	4544.0	4544.0
gl571.1	4164.0	4060.0	4002.0
gl57	4033.0	3913.0	3906.5
gl583	4261.0	4320.0	4320.0
gl707	4400.0	4364.0	4125.0
gl710	4160.0	4130.0	4165.0
gl747.3	4186.0	4170.0	4337.0
gl757	4008.0	4100.0	4259.0
gl782	4400.0	4590.0	4330.0
gl798	4300.0	4250.0	4015.5
gl801	4350.0	4060.0	4060.0
gl818	4462.5	4444.0	4430.0
gl826.1	4287.5	4350.0	4350.0
gl830	4269.0	4100.0	4065.5
gl842	3776.0	3720.0	3649.0
gl847a	4160.0	3470.0	3470.0
gl855	3831.0	3771.0	3771.0
gl884	4147.0	4130.0	4009.5
gl889a	4371.0	4350.0	4251.0
gl894	4270.0	4350.0	4350.0
gl898	4440.0	4350.0	4101.0
hip103150	4043.0	3900.0	3900.0
hip110245	4500.0	4416.0	4416.0
hip110714	4150.0	4060.0	4172.0
hip19410	4320.0	4238.0	4432.0
hip21865	4445.5	4432.0	4555.0
hip42108	4147.0	4343.0	4380.0
hip42910	4550.0	4373.0	4368.0
hip51073	4367.0	4310.0	4305.0
hip51263	4343.0	4350.0	4403.5
hip56838	4450.0	4060.0	4060.0
hip78395	4500.0	4414.0	4414.0
hip80083	4200.0	4800.0	4800.0

Notes. Hereafter $T_{\text{eff}}^{(\text{clos.})}$ was defined as the value found in VizieR catalogues closest to our inverted $T_{\text{eff}}^{(\text{inv.})}$, while $T_{\text{eff}}^{(\text{med.})}$ is the median of catalogues values.

Finally, to characterize our results, we first removed the eight above mentioned outliers from our objects list, which is about 14% of the original sample. Considering reference values as the one closest to our inverted effective temperature, we obtain a (mean signed difference or) bias of 21 K and a standard deviation of 90 K. If we use the median value as reference, the bias is 60 K and the standard deviation 132 K.

5. Conclusion

We have shown that the PCA-based inversion method of Paletou et al. (2015) provides realistic values for the effective temperature of late-type stars. Comparisons made between our estimates and effective temperature data found in the available literature reveals the existence of some strong discrepancies for a few objects. These discrepancies are most often related to very limited samples of estimates, so that additional investigations are clearly required for these objects.

These should consist in using different synthetic spectra, which can produce other radiative modelling tools such as Marcs (Gustafsson et al. 2008) or Phoenix, for cool stars (see e.g. Husser et al. 2013). The consideration of other spectral domains, and eventually the use of a combination of several distinct spectral domains should be explored too.

Our study also raises the more general question of the consistency between published (and made available) data, as well as the consistency between data provided by VizieR and Simbad services.

Acknowledgements. This research has made use of the VizieR catalogue access tool, CDS, Strasbourg, France. The original description of the VizieR service was published in A&AS 143, 23. This research has made use of the SIMBAD database, operated at CDS, Strasbourg, France.

References

- Adibekyan, V. Zh., Sousa, S. G., Santos, N. C., et al. 2012, *A&A*, **545**, A32
- Allende Prieto, C., Barklem, P. S., Lambert, D. L., & Cunha, K. 2004, *A&A*, **420**, 183
- Ammons, S. M., Robinson, S. E., Strader, J. et al. 2006, *ApJ*, **638**, 1004
- Carretta, E. 2013, *A&A*, **557**, A128
- Gaidos, E., Mann, A. W., Lépine, S., et al. 2014, *MNRAS*, **443**, 2561
- Gray, R. O., Corbally, C. J., Garrison, R. F., et al. 2006, *AJ*, **132**, 161
- Gustafsson, B., Edvardsson, B., Eriksson, K., et al. 2008, *A&A*, **486**, 951
- Hubeny, I., & Lanz, T. 1992, *A&A*, **262**, 501
- Husser, T.-O., Wende-von Berg, S., Dreizler, S., et al. 2013, *A&A*, **553**, A6
- Kordopatis, G., Gilmore, G., Steinmetz, M., et al. 2013, *AJ*, **146**, 134
- Kurucz, R. L. 1992, *Rev. Mex. Astron. Astrofis.*, **23**, 45
- Kurucz, R. L. 2005, *Mem. Soc. Astron. It. Supp.*, **8**, 14
- Lafrasse, S., Mella, G., Bonneau, D., et al. 2010, *SPIE Conf. Astronomical Instrumentation*, **77344E**, 140
- McDonald, I., Zijlstra, A. A., & Boyer, M. L. 2012, *MNRAS*, **427**, 343
- Morales, J. C., Ribas, I., & Jordi, C. 2008, *A&A*, **478**, 507
- Paletou, F., & Zolotukhin, I. 2014, ArXiv e-prints [arXiv:1408.7026]
- Paletou, F., Böhm, T., Watson, V., & Trouilhet, J.-F. 2015, *A&A*, **573**, A67
- Petit, P., Louge, T., Théado, S., et al. 2014, *PASP*, **126**, 469
- Prugniel, P., Soubiran, C., Koleva, M., & Le Borgne, D. 2007, ArXiv e-prints [arXiv:astro-ph/0703658]
- Santos, N. C., Sousa, S. G., Mortier, A., et al. 2013, *A&A*, **556**, A150
- Sousa, S. G., Santos, N. C., Israelian, G., Mayor, M., & Udry, S. 2011, *A&A*, **533**, A141
- van Leeuwen, F. 2007, *A&A*, **474**, 653
- Wright, C. O., Egan, M. P., Kraemer, K. E., & Price, S. D. 2003, *AJ*, **125**, 359

A new method for the inversion of atmospheric parameters of A/Am stars

Contexte : Nous présentons une procédure automatisée qui estime simultanément la température effective T_{eff} , la gravité de surface $\log(g)$, la métallicité $[Fe/H]$ et la vitesse de rotation équatoriale projetée $v \sin(i)$ pour des étoiles A et Am “normales”. La procédure est basée sur la méthode d’inversion utilisant l’analyse en composantes principales que nous avons publiée dans un article récent.

But : Un échantillon de 322 spectres de haute résolution pour des étoiles de type F0 à F9, récupérées des bases Polarbase, SOPHIE et ELODIE a été utilisé pour tester cette technique sur des données réelles. Nous avons sélectionné la région allant de 4 400 Å à 5 000 Å car elle contient beaucoup de raies métalliques et la raie de Balmer H_β .

Méthodes : Utilisant trois échantillons de données à des résolutions $R = 42\,000$, $65\,000$ et $76\,000$, approximativement 6.6×10^6 spectres synthétiques ont été calculés pour construire une grande base d’apprentissage. L’algorithme de “power iteration” a été utilisé pour calculer les composantes principales (PC). La projection des spectres sur ces PC permet d’avoir une métrique de comparaison efficace dans un espace de faible dimension. Les spectres bien connus des étoiles de type A0 et A1, Vega et Sirius A, ont été utilisés comme contrôle des trois bases de données. Les spectres d’autres étoiles de type A bien connues ont été utilisés pour caractériser la précision de la méthode d’inversion.

Résultats : Nous avons inversé tous les spectres observés et estimé les paramètres atmosphériques. Après le retrait de quelques “outliers”, la méthode basée sur l’ACP a été très efficace pour l’inversion de T_{eff} , $[Fe/H]$ et $v \sin(i)$ des étoiles de type A/Am. Les paramètres estimés sont en accords avec les déterminations antérieures. L’utilisation d’une approche statistique a établie des déviations de 150 K, 0.35 dex, 0.15 dex, 2 km/s pour T_{eff} , $\log(g)$, $[Fe/H]$ et $v \sin(i)$ par rapport aux valeurs de la littérature pour les étoiles de type A.

Conclusions : L’inversion basée sur l’ACP montre qu’elle est un outil très rapide, pratique et fiable pour l’estimation des paramètres stellaires des étoiles FGK et A, mais aussi pour l’estimation des températures effectives des étoiles de type M.

A new method for the inversion of atmospheric parameters of A/Am stars^{★,★★}

M. Gebran¹, W. Farah¹, F. Paletou^{2,3}, R. Monier^{4,5}, and V. Watson^{2,3}

¹ Department of Physics and Astronomy, Notre Dame University-Louaize, PO Box 72, Zouk Mikael, Lebanon
 e-mail: mgebran@ndu.edu.lb

² Université de Toulouse, UPS-Observatoire Midi-Pyrénées, IRAP, 31000 Toulouse, France

³ CNRS, Institut de Recherche en Astrophysique et Planétologie, 14 av. E. Belin, 31400 Toulouse, France

⁴ LESIA, UMR 8109, Observatoire de Paris, place J. Janssen, Meudon, France

⁵ Laboratoire Lagange, Université de Nice Sophia, Parc Valrose, 06100 Nice, France

Received 28 December 2015 / Accepted 2 March 2016

ABSTRACT

Context. We present an automated procedure that simultaneously derives the effective temperature T_{eff} , surface gravity $\log g$, metallicity $[\text{Fe}/\text{H}]$, and equatorial projected rotational velocity $v \sin i$ for “normal” A and Am stars. The procedure is based on the principal component analysis (PCA) inversion method, which we published in a recent paper.

Aims. A sample of 322 high-resolution spectra of F0-B9 stars, retrieved from the Polarbase, SOPHIE, and ELODIE databases, were used to test this technique with real data. We selected the spectral region from 4400–5000 Å as it contains many metallic lines and the Balmer H β line.

Methods. Using three data sets at resolving powers of $R = 42\,000$, 65 000 and 76 000, about $\sim 6.6 \times 10^6$ synthetic spectra were calculated to build a large learning database. The online power iteration algorithm was applied to these learning data sets to estimate the principal components (PC). The projection of spectra onto the few PCs offered an efficient comparison metric in a low-dimensional space. The spectra of the well-known A0- and A1-type stars, Vega and Sirius A, were used as control spectra in the three databases. Spectra of other well-known A-type stars were also employed to characterize the accuracy of the inversion technique.

Results. We inverted all of the observational spectra and derived the atmospheric parameters. After removal of a few outliers, the PCA-inversion method appeared to be very efficient in determining T_{eff} , $[\text{Fe}/\text{H}]$, and $v \sin i$ for A/Am stars. The derived parameters agree very well with previous determinations. Using a statistical approach, deviations of around 150 K, 0.35 dex, 0.15 dex, and 2 km s⁻¹ were found for T_{eff} , $\log g$, $[\text{Fe}/\text{H}]$, and $v \sin i$ with respect to literature values for A-type stars.

Conclusions. The PCA inversion proves to be a very fast, practical, and reliable tool for estimating stellar parameters of FGK and A stars and for deriving effective temperatures of M stars.

Key words. stars: fundamental parameters – stars: early-type – methods: numerical

1. Introduction

Sky and ground-based surveys can produce hundreds of terabytes (TB) and up to 100 (or more) petabytes (PB), both in the image data archive and in the object catalogues (Borne 2013). For instance, the ESA space mission *Gaia* (Perryman et al. 2001), launched in December 2013, is expected to produce a final archive of about 1 PB (10¹⁵ bytes) through five years of exploration. The Large Synoptic Survey Telescope (LSST) project will conduct a ten-year survey of the sky that will deliver about 15 terabytes (TB) of raw data per night (Kantor et al. 2007). The total amount of data collected over the ten years of operation will be 60 PB, and processing these data will produce a 15 PB catalogue database. The need for methods that can handle and analyse this large amount of information is evident.

Principal component analysis (PCA) has been used to invert stellar spectra to determine stellar fundamental parameters of stars (Bailer-Jones et al. 1998; Re Fiorentin et al. 2007; Paletou et al. 2015b,c). This technique has proven to be a fast

and reliable tool to invert observed high-resolution spectra. Unlike classical techniques such as least-squares fitting, PCA is based on a dimensionality reduction for the fast search of the nearest neighbour of an observed spectrum in a learning database. This technique can become obviously advantageous when dealing with large number of extended spectra collected at various spectral resolutions during sky surveys such as the Sloan Digital Sky Survey (SDSS; York et al. 2000), the RAdial Velocity Experiment (RAVE; Steinmetz 2003; Munari et al. 2005), and more recently the *Gaia* ESO Survey (GES; Gilmore et al. 2012).

A description of the PCA-based inversion technique for stellar parameters that we use can be found in Paletou et al. (2015b), where the authors applied this technique to derive fundamental parameters from high-resolution spectra of FGK stars. Paletou et al. (2015c) have shown that the range of application of the principal component analysis-based inversion method can be extended to dwarf M stars to derive effective temperatures.

We extend the work of Paletou et al. (2015b) to A-type stars whose effective temperatures range between 7000 and 10 000 K. A large variety of physical processes are at play in the envelopes of A/Am stars. Chemical peculiarity and pulsation are observed among main-sequence A and F stars. Chemical peculiarity is a signature of the occurrence of transport processes

* Based on data retrieved from the Polarbase, SOPHIE, and ELODIE archives.

** Table 2 is only available at the CDS via anonymous ftp to cdsarc.u-strasbg.fr (130.79.128.5) or via <http://cdsarc.u-strasbg.fr/viz-bin/qcat?J/A+A/589/A83>

competing with radiative diffusion (Zahn 2005; Richer et al. 2000; Richard et al. 2001, 2002; Vick et al. 2010). Chemical peculiarity exist in different ways in A-type stars (Am, Ap, λ Bootis stars...). These peculiarities set constraints on self-consistent evolutionary models of these objects, including various particle transport processes (Gebran et al. 2008; Gebran & Monier 2008; Gebran et al. 2010). The stellar atmospheric parameters are the prerequisites to any detailed abundance analysis, and for that reason, these parameters should be derived with good accuracy.

To do so, we calculated around 6.6×10^6 synthetic spectra and used them as three learning databases. These synthetic spectra were compared to a collection of observed spectra to find the best global fit between the two sets. The A-type stars spectra that we analysed were selected from the PolarBase¹ (Petit et al. 2014) stellar library at two different resolutions ($R \sim 65\,000$ and $R \sim 76\,000$), and from the SOPHIE² and ELODIE³ archives (Moultaka et al. 2004) ($R \sim 75\,000$ and $42\,000$, respectively). We restricted our study to the wavelength range $4400\text{--}5000\text{ \AA}$. In this region, free from telluric lines, we can find many metallic lines and mostly many iron transitions (Fe I and Fe II) with accurate atomic data. The Balmer line $H\beta$ around $\sim 4860\text{ \AA}$ is also present in this spectral range. As mentioned by Smalley (2005), the advantage of the global spectral fitting technique is that it can be automated for a very large number of stellar observations. It can also be used for low-resolution spectra or spectra of fast rotators, where line blending is important ($v \sin i$ can go up to 300 km s^{-1} for A-type stars). Dealing with large multi-dimensional grid of synthetic spectra, the combination of the metallic and Balmer lines is essential for the derivation of the effective temperature T_{eff} , surface gravity $\log g$, metallicity $[\text{Fe}/\text{H}]$, and projected equatorial rotational velocity $v \sin i$ (Smalley 2005).

As a result of their sensitivity to effective temperature and surface gravity, hydrogen line profiles should be included in all automatic procedures based on spectral fitting to obtain realistic atmospheric parameters for stars hotter than 5500 K (Ryabchikova et al. 2015). The Balmer lines in B to F stars are in general well matched by LTE computations, and they are practically identical to the non-LTE predictions (Przybilla & Butler 2004). The region around $H\beta$ was chosen because it harbours several iron lines with accurate atomic parameters. The Balmer lines provide an excellent diagnostic of T_{eff} for stars cooler than $\sim 8000\text{ K}$ (Gray 1992; Heiter et al. 2002). For stars hotter than 8000 K , Balmer line profiles, in particular their wings, are sensitive to both temperature and gravity (Gray 1992; Smalley 2005). Metal line diagnostics are also good indicators of T_{eff} , $\log g$, and $[\text{Fe}/\text{H}]$ as the best-fit synthetic spectrum corresponds to the spectrum that yields the same abundance for different ionization stages of the same element (ionization balance) and in the same abundance for all excitation potentials of the same element in a given ionization stage (excitation potential). On the other hand, lines such as the Mg II triplet around $4481\text{ \AA}$ and the Fe II lines between 4500 and $4530\text{ \AA}$ are excellent indicators of $v \sin i$ (see for example Royer et al. 2014). For all the above mentioned reasons, we applied the global spectral fitting technique to the whole spectral range from 4400 and $5000\text{ \AA}$. The derived parameters are obviously model dependent, but the models that we are using (Sect. 2.2) are well suited to describe the atmospheres and spectra of “normal” A and Am stars.

The selection of the observed spectra and construction of the learning databases are detailed in Sect. 2. The PCA and the derivation of the eigenvectors and coefficients are recalled in Sect. 3. Radial velocity correction and re-normalization of the spectra are discussed in Sect. 4. The results of the inversion are presented in Sect. 5. A few outliers have been analysed in detail in Sect. 6. Discussion and conclusion are gathered in Sect. 7.

2. Observations and learning database

2.1. Target selection

The PolarBase contains reduced stellar spectra collected with the NARVAL and ESPaDOnS high-resolution spectropolarimeters. These two spectropolarimeters are mounted on the 2 m aperture *Télescope Bernard Lyot* (TBL) in France and on the 3.6 m aperture CFHT telescope in Hawaii, respectively. Both instruments have a spectral resolving power of $65\,000$ in polarimetric mode and $76\,000$ when used for classical spectroscopy. As a first step, we have selected all high-resolution Echelle spectra of stars whose spectral type ranges from F0 up to B9 from the database at two different resolutions and in spectroscopy mode only. Fifty spectra were thus retrieved.

We have also retrieved A-type stars spectra from the ELODIE and SOPHIE archives. These archives are on-line databases of high-resolution Echelle spectra and radial velocities. The ELODIE is a fiber-fed, cross-dispersed Echelle spectrograph that was attached to the 1.93-m telescope at Observatoire de Haute-Provence (OHP; Baranne et al. 1996). The ELODIE recorded in a single exposure a spectrum extending from $3850\text{ \AA}$ to $6811\text{ \AA}$ at a resolving power of about $42\,000$ on a relatively small CCD (1024×1024). SOPHIE is the Echelle spectrograph mounted on the 1.93-m telescope at the OHP since September 2006. SOPHIE spectra stretch from 3870 to $6940\text{ \AA}$ in 39 orders with two different spectral resolutions: the high-resolution mode (HR; $R \sim 75\,000$) and the high efficiency mode (HE; $R = 39\,000$). All the observations that we used were observed in the HR mode at the same resolution of NARVAL spectra. Spectra of 279 stars, with signal-to-noise ratio larger than 100, were downloaded from the SOPHIE and ELODIE archives. Overall, we selected 322 stars for our study. The identifications of the selected stars are collected in Table 2. These data correspond to PolarBase, ELODIE, and SOPHIE spectra, respectively. Inverted effective temperatures, surface gravities, projected equatorial rotational velocities, metallicities, and radial velocities are collected in this table (see Sect. 5).

2.2. Learning database

The learning database was constructed using synthetic spectra following Paletou et al. (2015c). Three sets of grids were used, at a resolution of $65\,000$, $76\,000$, and $42\,000$, corresponding to the resolution of Polarbase, ELODIE and SOPHIE spectra, respectively. These three grids are identical in terms of the parameters (T_{eff} , $\log g$, $[\text{Fe}/\text{H}]$, $v \sin i$, ξ_i), and only the resolution differs. Line-blanketed ATLAS9 model atmospheres (Kurucz 1992) were calculated for the purpose of this work. ATLAS9 models are LTE plane parallel and assume radiative and hydrostatic equilibrium. This ATLAS9 version uses the opacity distribution function (ODF) of Castelli & Kurucz (2003). We included convection in the atmospheres of stars cooler than 8500 K . Convection was treated using a mixing length parameter of 0.5 for $7000\text{ K} \leq T_{\text{eff}} \leq 8500\text{ K}$, and 1.25 for $T_{\text{eff}} \leq 7000\text{ K}$, following Smalley’s (2004) prescriptions.

¹ <http://polarbase.irap.omp.eu>

² <http://atlas.obs-hp.fr/sophie/>

³ <http://atlas.obs-hp.fr/elodie/>

The grid of synthetic spectra was computed using SYNSPEC48 (Hubeny & Lanz 1992). The effective temperatures of the model atmospheres range from 6800 and 11 000 K with a step of 100 K, and logarithm of surface gravities between 2.0 and 5.0 with a step of 0.1 dex, respectively. The projected equatorial rotational velocities $v \sin i$ were varied from 0 up to 300 km s⁻¹ with a non-constant step (2, 5, and 10 km s⁻¹), and the metallicity was scaled from -2.0 dex up to +2.0 dex with respect to the Grevesse & Sauval (1998) solar value with a step of 0.1 dex. The synthetic spectra were computed from 4400 Å up to 5000 Å with a wavelength step of 0.05 Å. This range contains many moderate and weak metallic lines in different ionization stages. These weak metallic lines (indicators for $v \sin i$ and [Fe/H]) and the Balmer line (indicator for T_{eff} and $\log g$) are not sensitive to small changes of the microturbulent velocity. For that reason, we fixed the microturbulent velocity (ξ_r) to 2.0 km s⁻¹ in all models. This value corresponds to the average microturbulent velocity in A-type stars (Gebran et al. 2014). The line list used in the synthetic spectra calculation was constructed from Kurucz ghyperall.dat⁴ and modified with more recent and accurate atomic data retrieved from the VALD⁵ and the NIST⁶ databases. At each resolution, we ended up with ~2.2 million spectra as a learning database.

3. Principal component analysis

Our usage of the PCA for the inversion process is influenced by the work of Rees et al. (2000) and Paletou et al. (2015b,c). As mentioned in Sect. 1, the PCA is a reduction of dimensionality technique that retains the variation present in the data set (Jolliffe 1986) as much as possible. It searches for basis vectors that represent most of the variance in a given database. The reduction of dimensionality allowed by the PCA is directly used to build a specific metric from which a nearest-neighbour(s) search is made between an observed data set (observed spectra in the present work) and the content of a learning database (synthetic spectra in the present work). The learning database could be constructed with a set of observed spectra with well-known parameters (Paletou et al. 2015b). The PCA becomes mainly advantageous when dealing with high- and low-resolution spectra that cover a very broad bandwidth as the reduction of dimensionality allows a very fast processing of the data. In what follows, we describe the characteristics of this technique in the context of our study.

3.1. Nearest neighbour search

Representing the learning database in a matrix $\mathbf{S}$ of size $N_{\text{spectra}} = 2.2 \times 10^6$ by $N_\lambda = 12\,000$, $\bar{\mathbf{S}}$ is the average of $\mathbf{S}$ along the N_{spectra} -axis. Then, we derive the eigenvectors $\mathbf{e}_k(\lambda)$ of the variance-covariance matrix $\mathbf{C}$ defined as

$$\mathbf{C} = (\mathbf{S} - \bar{\mathbf{S}})^T \cdot (\mathbf{S} - \bar{\mathbf{S}}) \quad (1)$$

with a dimension of $N_\lambda \times N_\lambda$. These eigenvectors of the symmetric $\mathbf{C}$ matrix are then sorted in decreasing eigenvalues magnitude; they are more usually referred to as “principal components”. Each spectrum of the data set $\mathbf{S}$ is represented by a small number of coefficients p_{jk} defined as

$$p_{jk} = (\mathbf{S}_j - \bar{\mathbf{S}}) \cdot \mathbf{e}_k. \quad (2)$$

⁴ <http://kurucz.harvard.edu>

⁵ <http://www.astro.uu.se/~vald/php/vald.php>

⁶ <http://physics.nist.gov>

The main point of reducing the dimensionality is the selection of $k_{\text{max}} \ll N_\lambda$. We show in Sect. 3.2 that $k_{\text{max}} = 12$ is a reasonable choice.

To follow the same notation as in Paletou et al. (2015b), we denote the observed spectrum by $O(\lambda)$ having the same wavelength range and sampling as the learning data set spectra. At this stage, the observations should be corrected for radial velocity (details in Sect. 4). We compute the set of projection coefficients Q_k defined as

$$Q_k = (O - \bar{\mathbf{S}}) \cdot \mathbf{e}_k. \quad (3)$$

The nearest neighbour is then found by minimizing a χ^2 in the low-dimensional space of the coefficients

$$d_j^{(O)} = \sum_{k=1}^{k_{\text{max}}} (Q_k - p_{jk})^2, \quad (4)$$

where j covers the number of synthetic spectra. The parameters of the synthetic spectrum with the minimum d are considered for the inversion.

3.2. Power iteration method

The derivation of the leading eigenvectors of $\mathbf{C}$ is the most critical step when performing PCA, and it becomes increasingly time consuming as the size of learning data set increases. Furthermore, techniques, such as singular value decomposition (SVD) that were used in previous studies, become inapplicable in this case as a result of computer memory shortage. In our particular case, at each of the three resolutions, the files containing the matrix $\mathbf{S}$ are ~240 GB in size. To bypass this problem, we used an algorithm that is applicable independent of the size of the data set. We implemented the online power iteration algorithm (also called Von Mises iteration; Demmel 1997) to estimate the k_{max} principal components of the data set $\mathbf{S}$. This algorithm reads and operates on synthetic spectra of the learning data set one by one, avoiding loading all the data at once. The procedure, moreover, does not need any typical matrix factorization such as SVD, and is efficient when applied to large data sets (Roweis 1998). The procedure consists of the following.

First we find the eigenvector $\mathbf{e}_1$ with a maximal eigenvalue λ_1 (i.e. the first principal component), by iterating

$$x_{n+1} = \frac{\mathbf{C}x_n}{\|\mathbf{C}x_n\|}, \quad (5)$$

with an initial random non-zero vector x_0 . Practically $\mathbf{C}x_n$ cannot be computed using standard technique as one should load the full data set $\mathbf{S}$ at once to calculate $\mathbf{C}$. To overcome this problem, we evaluate “ $\mathbf{C}x_n$ ” for each iteration on x_n while operating on each spectrum at a time by performing

$$\begin{aligned} \mathbf{C}x_n &= (\mathbf{S} - \bar{\mathbf{S}})^T (\mathbf{S} - \bar{\mathbf{S}}) \cdot x_n \\ &= \sum_{i=1}^{N_{\text{spectra}}} ((S_i - \bar{\mathbf{S}}) \cdot x_n) \cdot (S_i - \bar{\mathbf{S}}), \end{aligned} \quad (6)$$

where S_i is the i th spectrum in the data set $\mathbf{S}$. The x_n iterations converge to the eigenvector $\mathbf{e}_1$ that corresponds to the largest eigenvalue.

Once we find a good approximation for $\mathbf{e}_1$, we perform the Gram-Schmidt process to remove the contribution of $\mathbf{e}_1$ from $\mathbf{S}$ by computing

$$\mathbf{S}_1 = \mathbf{S} - [(\mathbf{S} - \bar{\mathbf{S}}) \cdot \mathbf{e}_1] \mathbf{e}_1. \quad (7)$$

Again, the above should be performed iteratively over each spectrum. Next, to compute e_2 , the power method (Eq. (5)) is performed on

$$\mathbf{C}_1 = (\mathbf{S}_1 - \bar{\mathbf{S}})^T \cdot (\mathbf{S}_1 - \bar{\mathbf{S}}). \quad (8)$$

Once e_2 is estimated, the Gram-Schmidt process is applied again to yield

$$\mathbf{S}_2 = \mathbf{S}_1 - [(\mathbf{S}_1 - \bar{\mathbf{S}}) \cdot e_2]e_2. \quad (9)$$

And so on, until we compute all the needed eigenvectors. To determine the optimal number of iterations needed for the derivation of the eigenvectors, we constructed several smaller databases and derived their respective eigenvectors using the SVD technique and the power iteration. For each database, the derived eigenvectors from both techniques were compared. Using power iteration, 25 iterations were required to reach a relative difference between eigenvectors smaller than 0.1%. The use of power iteration allowed the increase of the number of spectra in the learning database by a factor of 10 (200 000 in case of the SVD to 2 000 000 in case of power iteration).

Only the first 12 eigenvectors were considered and that was justified by the calculation of the reconstruction error of the spectra defined as

$$E(k_{\max}) = \left(\left(\frac{|\bar{\mathbf{S}} + \sum_{k=1}^{k_{\max}} p_{jk} \mathbf{e}_k - \mathbf{S}_j|}{S_j} \right) \right). \quad (10)$$

Using only 12 eigenvectors, the average reconstruction error is found to be smaller than 1%, a result similar to that derived in [Paletou et al. \(2015b\)](#).

4. Preparation of the observed spectra

In order to adjust a synthetic spectrum to an observed spectrum, lines of both spectra should be aligned, as accurately as possible, on the same wavelength scale and both fluxes should be at the same level. The flux matching is achieved by rectifying the observed and synthetic spectra to their local continua.

4.1. Radial velocity correction

A wavelength shift between observations and synthetic spectra could drastically affect the derived parameters of the first neighbour. [Paletou et al. \(2015b\)](#) showed that the radial velocity v_r should be known to an accuracy of $c/4R$, where c is the speed of light in vacuum. This requirement translates into a velocity accuracy of $\sim 1 \text{ km s}^{-1}$ needed at a spectral resolution of 76 000. We corrected all the spectra from radial velocity (given in the last column of Table 2) before the inversion procedure. The radial velocities were determined with the classical cross-correlation technique originally described by [Tonry & Davis \(1979\)](#). The observed spectra were cross correlated using a synthetic template, in the same wavelength range, corresponding to the parameters $T_{\text{eff}} = 8500 \text{ K}$, $\log g = 4.0 \text{ dex}$, $[\text{Fe}/\text{H}] = 0$, and $v_e \sin i = 0 \text{ km s}^{-1}$. Many of the Fe II, Cr II, Ti II, and Mg II atomic lines are mainly present in all sub-types of A-type stars (from A0 to A9). The wavelength of the Fe II lines in the 4500–4550 Å region are accurately measured to within $\sim 0.01 \text{ Å}$ according to the NIST database. For that reason, the selected template corresponds to a typical A5V type star.

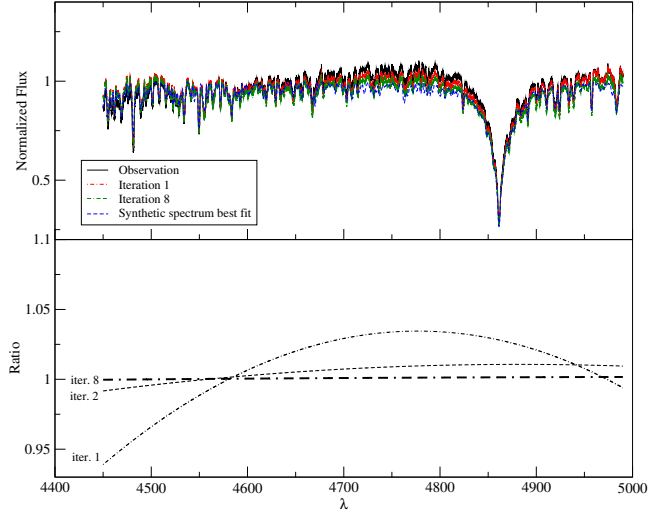


Fig. 1. *Upper panel:* continuum correction using eight iterations for the inversion of HD 97230. Only iterations 1 and 8 are shown for the sake of clarity. The initial observation is indicated with a solid line, the 1st iteration is indicated with dash-dotted line, the 8th iteration is indicated with double dash-dotted line, and the final best-fit synthetic spectrum is indicated with a dashed line. *Lower panel:* ratios between the re-normalized spectrum of the observation and the first neighbour best-fit synthetic spectrum for iterations 1, 2, and 8.

4.2. Renormalization of the flux

On the other hand, we also noticed that the flux normalization procedure is critical in this work and that an improper normalization could affect drastically the derived parameters. We proceeded using [Gazzano et al. \(2010\)](#)'s procedure to correct for this effect. Each originally normalized observed spectrum was divided by the synthetic spectrum calculated with the first inverted atmospheric parameters (T_{eff} , $\log g$, $[\text{Fe}/\text{H}]$, and $v \sin i$). The residuals were then cleaned from lines with a sigma clipping, rejecting points above or below 1σ . The remaining points were fitted with a second degree polynomial. The observation was then divided by this polynomial in order to obtain a new normalized spectrum. This new spectrum was inverted again to find a new set of parameters. This procedure was repeated until the inverted parameters remained unchanged from one iteration to the next. The number of iterations depends on the original shape of the normalized observed spectrum. On average, ten iterations were sufficient to reach convergence. In Fig. 1, we show the application of the Gazzano et al. (2010) procedure to the inversion of the spectrum of HD 97230; only eight iterations were required in this example.

5. A stars spectra inversions

5.1. Vega and Sirius A

The first tests were carried out using the two well-studied A-type stars, Vega (HD 172167, A0V) and Sirius A (HD 48915, A1m). The spectra of Vega were retrieved from the PolarBase and SOPHIE archive. The ESPaDOs spectrum was inverted using the learning database at a resolution of 65 000, whereas the SOPHIE spectrum was inverted using the database at a resolution of 76 000. Both inversions yielded the same results for all parameters. We found for $\{T_{\text{eff}}, \log g, [\text{Fe}/\text{H}], v \sin i\}$ values of $\{9500 \text{ K}, 4.0 \text{ dex}, -0.5 \text{ dex}, 25 \text{ km s}^{-1}\}$, whereas the medians of

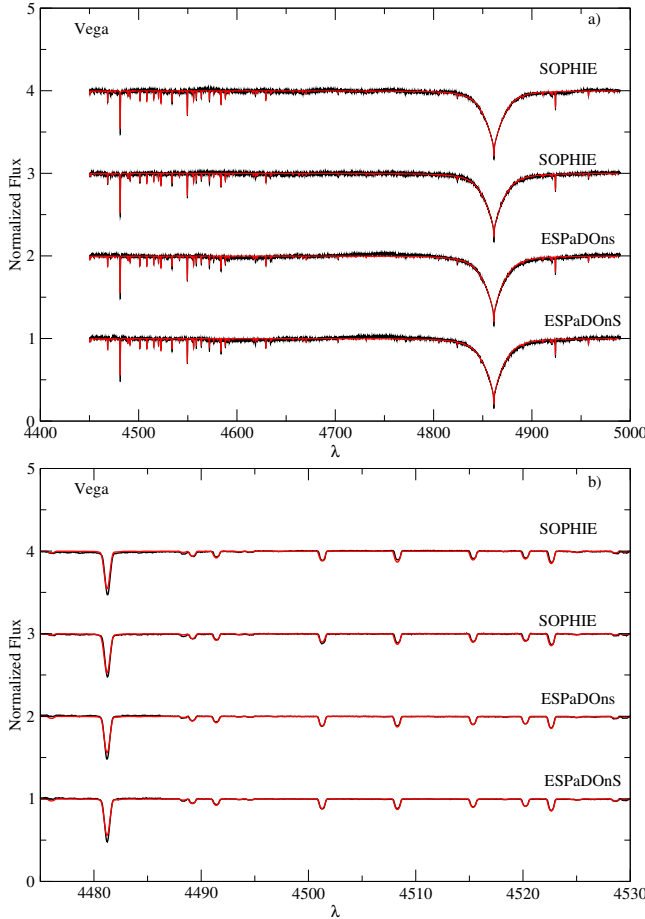


Fig. 2. First neighbour fit of synthetic spectra to the observed spectra for the four spectra of Vega, observed at different epochs, retrieved from PolarBase and SOPHIE databases. Observed spectra are shown in black and the synthetic spectra are shown in red. Part **a)** represents the overall spectrum, whereas part **b)** indicates the region sensitive to $[\text{Fe}/\text{H}]$ and $v \sin i$. The wavelengths are in Å. Spectra were offset by arbitrary amounts.

the parameters found by querying VizieR⁷ are {9520 K, 3.98 dex, -0.53 dex, 23 km s^{-1} }. The inverted parameters of Vega agree very well with all previous determinations, such as Takeda et al. (2008), Hill et al. (2010), and the fundamental effective temperature derived by Code et al. (1976) from the integrated flux and the angular diameter. Figure 2a shows the overall adjustment of four observed spectra of Vega with the first neighbour synthetic spectra. These spectra were observed at different epochs. The good agreement between the inverted metallicity and the projected equatorial rotational velocity of Vega with those from the catalogue can be seen in Fig. 2b, where we show the adjustment in the region sensitive to these two parameters. The spectrum of Sirius A was retrieved from the ELODIE database. This spectrum was inverted using the learning database at a resolution of 42 000. We found for $\{T_{\text{eff}}, \log g, [\text{Fe}/\text{H}], v \sin i\}$ values of {9800 K, 4.2 dex, 0.2 dex, 18 km s^{-1} }, whereas the medians of the parameters retrieved from a VizieR query are {9870 K, 4.30 dex, 0.36 dex, 15 km s^{-1} }. This result is in very good agreement with previous findings such as Hill & Landstreet (1993), Takeda et al. (2009), and Landstreet (2011).

⁷ <http://vizier.u-strasbg.fr/viz-bin/VizieR>

Table 2 shows the inverted parameters, the values that we retrieved from VizieR queries closest to our inverted values, and the median of queried catalogues values. All literature data were collected from VizieR catalogues using the ASTROQUERY⁸ Python modules (Paletou & Zolotukhin 2014). We kept the inverted parameters for all selected stars in this table even though many of these stars may be variables or binaries. Using Simbad queries, we listed, in the last column of Table 2, the peculiarity of the stars. Variables and binary systems were discarded in the characterization of the technique (Sect. 5.2).

The effective temperature is the only atmospheric parameter that can be retrieved from VizieR queries for all stars except for one exception. A small number of the selected stars do not have catalogue values for $\log g$, $v \sin i$, and $[\text{Fe}/\text{H}]$. Figure 3 shows the inverted effective temperature for each spectrum retrieved from Polarbase as well as the range in the effective temperatures retrieved from the catalogues (box-plots) and the median. The number of values found in all catalogues retrieved in VizieR for each star also appears in this figure. We found that the inverted parameters agree well with those retrieved in the various catalogues. This agreement also applies to A-type stars spectra retrieved from ELODIE and SOPHIE archives (Table 2). However, we found several outliers that are discussed in detail in Sect. 6. These outliers are depicted with a symbol “*” in Table 2. Figure 4 shows an example of an overall spectral fitting between the first neighbour synthetic spectrum and the observation for several A-type stars of our study.

5.2. Evaluation of the results and comparison with literature data

In order to evaluate the efficiency of our method, first we removed from the analysis all peculiar stars, such as Ae/Be, Ap, variable (δ Scuti, γ Doradus variables...), and binary stars (eclipsing binaries, spectroscopic binaries...) as the models that we are using do not take such peculiarities into account. ATLAS12 (Kurucz 2005) model atmospheres, dealing with opacity sampling (OS), can better reproduce the distribution of the thermodynamical quantities in the atmosphere of chemically peculiar stars. Peculiarities of Am stars are not strong enough to alter the temperature distribution in their atmospheres (Gebran et al. 2008; Catanzaro & Balona 2012). For that reason, we are confident that ATLAS9 can be used for moderate chemical peculiarity.

We ended up with 182 classified “normal” A or Am stars. To assess an automated and objective comparison with already published values that one may find in VizieR@CDS, for instance, we found it reasonable to adopt specific weights for the set of values we could retrieve for each object. These weights directly rely on the spread of catalogued values. For some of the stars, dispersions in the catalogued values were noticed, and were severe up to thousands of Kelvin for T_{eff} . As an example, HIP109119 has catalogued T_{eff} ranging between 5309 K and 8970 K. On the other hand, some stars had only one or no catalogued values at all for a particular parameter. This would definitely add a bias to the error calculation. For that reason we used a statistical approach that gives less weight to stars with a large spread in the catalogued values. For a star i and a parameter ψ (ψ could be one of the following parameters: T_{eff} , $\log g$, $[\text{Fe}/\text{H}]$, $v \sin i$), the difference between the inverted and median from the catalogued values of this parameter is weighted with the inverse of the interquartile range (IQR). IQR is defined as the difference between

⁸ astroquery.readthedocs.org

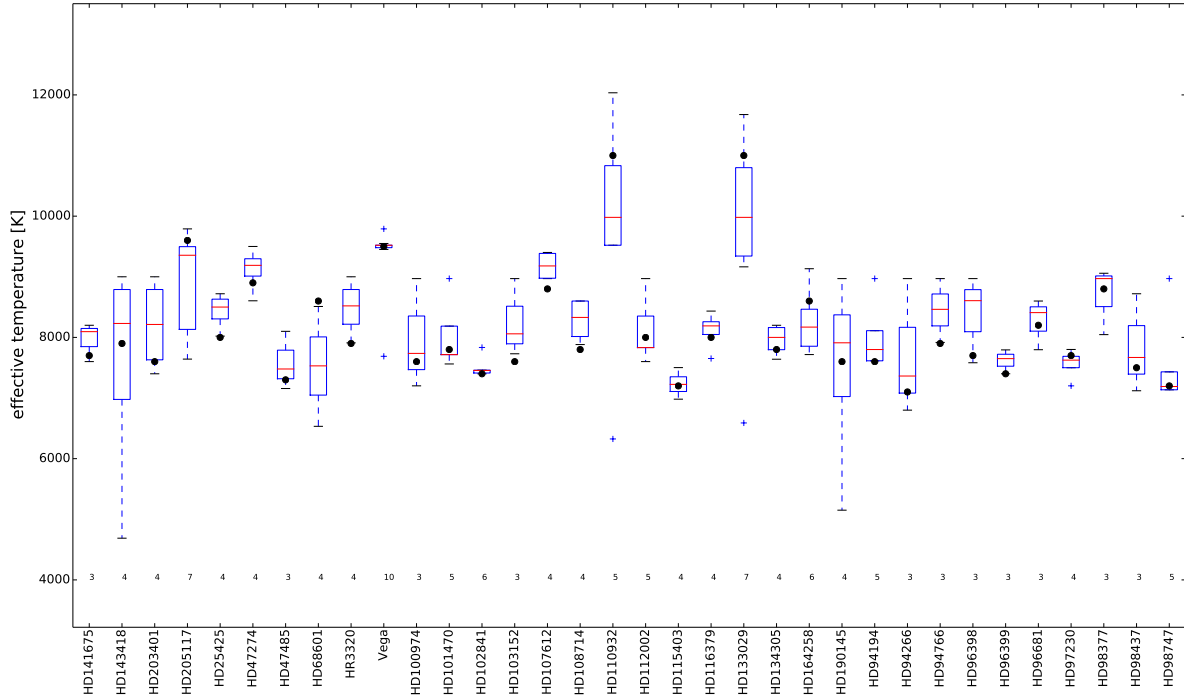


Fig. 3. Comparison between our estimate of effective temperatures (●), and the values we obtained from available VizieR catalogues, for PolarBase spectra. The latter collections are represented as classical boxplots. Objects we studied are listed along the horizontal axis. The number of values found among all VizieR catalogues are presented on a horizontal line around $T_{\text{eff}} \sim 4000$ K. The horizontal bar inside each box indicates the median (Q_2 value), while each box extends from first quartile, Q_1 , to third quartile Q_3 . Extreme values, still within a 1.5 times the interquartile range away from either Q_1 or Q_3 , are connected to the box with dashed lines. Outliers are denoted by a “+” symbol.

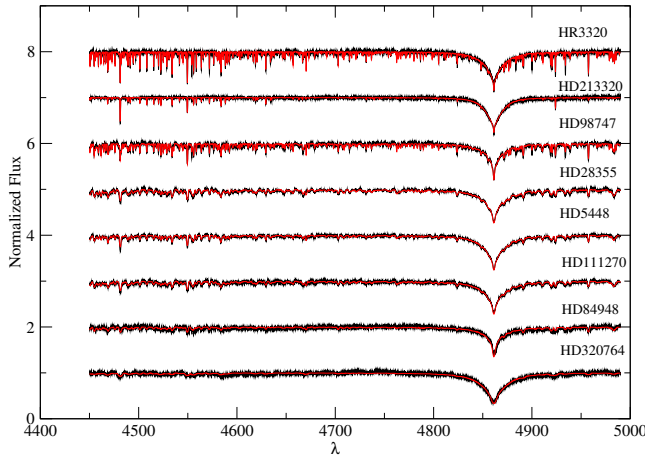


Fig. 4. First neighbour fit of synthetic spectra to the observed spectra. The wavelengths are in Å. Spectra were shifted up for clarity.

the third and first quartile of each set of values ($IQR = Q_3 - Q_1$). The deviation for the parameter ψ is derived following:

$$\Delta\psi_w = \frac{\sum_i w_i (\psi_i^{\text{inv}} - \psi_i^{\text{med}})}{\sum w_i}, \quad (11)$$

with w_i derived for each star and each parameter as follows:

$$w_i = \frac{1}{IQR_i}. \quad (12)$$

The associated standard deviation is found according to the following equation:

$$\sigma_w^2 = \frac{\sum_i w_i (\psi_i^{\text{inv}} - \psi_i^{\text{med}})^2}{\sum w_i}. \quad (13)$$

Stars with fewer than two catalogued values for a parameter ψ were omitted in the calculation of $\Delta\psi_w$ because of the zero value of their IQR . To estimate the parameters spread in the catalogued values for the remaining stars, we calculated the average IQR and the standard deviation. We found that, on average, the interquartile values for $\{T_{\text{eff}}, \log g, [\text{Fe}/\text{H}], v \sin i\}$ are $\{550 \text{ K}, 0.25 \text{ dex}, 0.17 \text{ dex}, 10.7 \text{ km s}^{-1}\}$ with a standard deviations of $\{1200 \text{ K}, 0.35 \text{ dex}, 0.23 \text{ dex}, 15.5 \text{ km s}^{-1}\}$. These numbers show that a large spread exists in the literature values, validating our approach in using a weighted mean. Applying Eq. (11), the derived deviations on the four parameters and their corresponding standard deviations are summarized in Cols. 2 and 3 of Table 1. Columns 4 to 11 of Table 1 show the absolute mean signed differences and standard deviations obtained by comparing our inverted parameters to the median and closest parameters in the catalogues from VizieR query for all stars (Cols. 4 to 7) and for the 19 reference stars (Cols. 8 to 11, see Sect. 5.2.1).

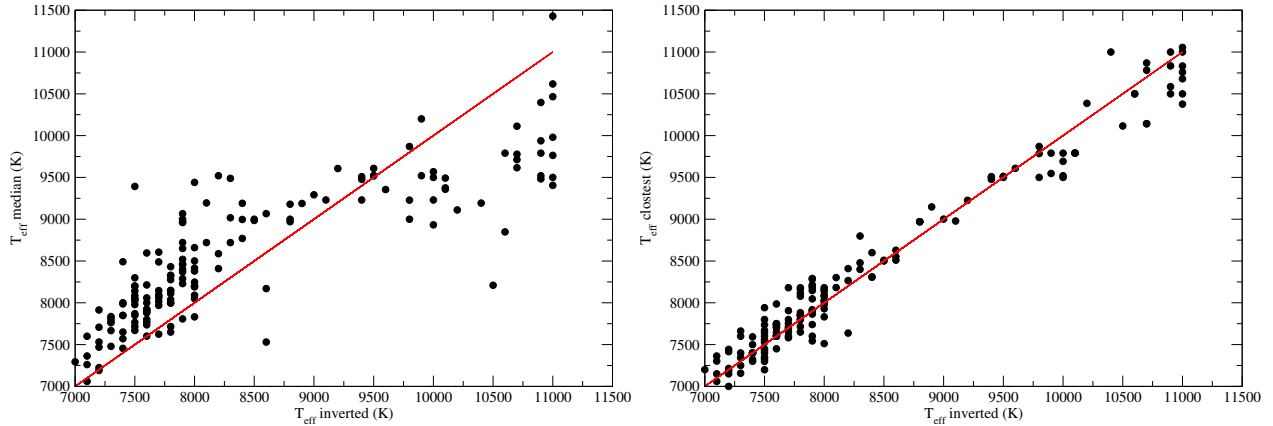
Effective temperature T_{eff}

Figure 5 shows the comparison between our inverted effective temperature and the catalogued effective temperatures for the 182 classified “normal” A or Am stars. The left panel of Fig. 5 represents the comparison with the median of the catalogued values, whereas the right panel shows the comparison with the closest. There is a systematic trend noticed in the left plot of Fig. 5. This is partly because of the large spread in the catalogued data. By comparing our inverted T_{eff} to the median value, we obtain an absolute mean signed difference of 110 K with a standard deviation of 630 K (Cols. 4 and 5 of Table 1). All stars (except HD 322676) have at least two values for catalogued T_{eff} . These catalogued effective temperatures are based on different bibliographical references using distinct tools/techniques

Table 1. Average weighted deviations between our inverted parameters and catalogued values (Δ_w) with their respective standard deviation (σ_w).

	Δ_w	σ_w	Δ_m	σ_m	Δ_c	σ_c	Δ_{19m}	σ_{19m}	Δ_{19c}	σ_{19c}
T_{eff} (K)	150	500	110	630	13	220	35	250	8.5	96
$\log g$ (dex)	0.35	0.30	0.40	0.60	0.40	0.30	0.18	0.40	0.17	0.30
[Fe/H] (dex)	0.15	0.25	0.15	0.25	0.12	0.20	0.09	0.15	0.08	0.13
$v \sin i$ (km s ⁻¹)	2.0	9.5	5.0	15.0	1.7	10.0	4.0	9.5	2.0	5.5

Notes. The Δ_m , σ_m , Δ_c , and σ_c show the unweighed average mean signed differences and their respective standard deviations between our inverted parameters and the median (m)/closest (c) of the catalogue values for the 182 classified normal A or Am stars and 19 reference (Δ_{19} and σ_{19}) stars.

**Fig. 5.** Left panel: comparison between our inverted effective temperatures and the median of the catalogued effective temperatures. Right panel: comparison between our inverted effective temperatures and the closest of the catalogued effective temperatures.

for their determination. Using spectroscopy, one should expect some differences between the derived effective temperatures and those determined with photometric techniques. The global spectral fitting that includes the Balmer lines provides an excellent T_{eff} diagnostic for stars cooler than ~ 8000 K because they do not depend on gravity (Gray 1992). It follows that one part of the differences between inverted and catalogued values is justified for stars with $T_{\text{eff}} > 8000$ K. For stars hotter than 8000 K, a degeneracy could exist between the T_{eff} and $\log g$. To overcome this degeneracy, we selected a wavelength range containing many metallic lines as described in Sect. 1. When comparing our inverted effective temperatures with the closest catalogued values, the correlation becomes clearer (correlation coefficient of 0.98). The absolute mean signed difference and standard deviation are found to be, in that case, 13 K and 220 K, respectively (Cols. 6 and 7 of Table 1). Taking into account the spread of catalogued data by applying Eq. (11), we found that on average our effective temperatures deviate about 150 K from the catalogued effective temperatures with a standard deviation of 500 K (Cols. 2 and 3 of Table 1).

Surface gravity $\log g$

As mentioned in Sect. 1, the H β line is not sensitive to gravity for $T_{\text{eff}} \leq 8000$ K, which applies to about 60% of our sample. For that reason, we expected that the global spectral fitting technique could derive incorrect $\log g$ values for these stars. The most reliable $\log g$ values for A-type stars are determined using photometric techniques, more precisely, from the *uvby* β calibration. The photometric values of $\log g$ are not affected by metallicity (Smalley & Dworetsky 1993). With a combination of spectroscopic and photometric techniques, the typical errors on the atmospheric parameters of a star is ± 100 K for T_{eff} and ± 0.2 dex for $\log g$ (Smalley 2005). The 0.35 dex weighted average deviation and its corresponding 0.30 dex standard deviation found

in the present work (Cols. 2 and 3 of Table 1), by applying Eq. (11), are justified by the use of only spectroscopic fitting technique and by the insensitivity of the H β line to gravity for $T_{\text{eff}} \leq 8000$ K. Comparing our inverted $\log g$ to the median of the catalogued $\log g$, we found an unweighed absolute mean signed difference of 0.40 dex with a standard deviation of about ~ 0.60 dex (Cols. 4 and 5 of Table 1).

Metallicity [Fe/H]

Catalogued values for metallicities exist for a small number of stars and most of them have only one determination of [Fe/H]. Statistically, our comparison with catalogued values is not reliable in this case, but the 0.15 dex that we derived using our statistical approach (Eq. (11)) is an average deviation for the derived metallicities. The standard deviation in that case is 0.25 dex (Eq. (13)). These values are identical to those derived using the unweighed differences (Cols. 4 and 5 of Table 1). Comparing with the closest catalogues values, these numbers are found to be 0.12 dex and 0.20 dex, respectively (Cols. 6 and 7 of Table 1).

Projected equatorial rotational velocity $v \sin i$

Most of the catalogued $v \sin i$ are based on the computation of the first zero of Fourier transform (FT) of Royer et al. (2002). Our spectroscopic $v \sin i$ agree very well with previous findings as shown in Fig. 6, where the comparison was carried out with the median and closest catalogued values. We found a correlation coefficient of 0.96 between our inverted values and the median values from Vizier catalogues with a regression coefficient (slope) of 0.98 ± 0.02 . The correlation coefficient reaches 0.988 when dealing with the closest values. By comparing one-to-one values between our inverted $v \sin i$ and the median of the catalogued values, we found an absolute mean signed difference of 5.0 km s⁻¹ and a standard deviation of 15.0 km s⁻¹ (Cols. 4

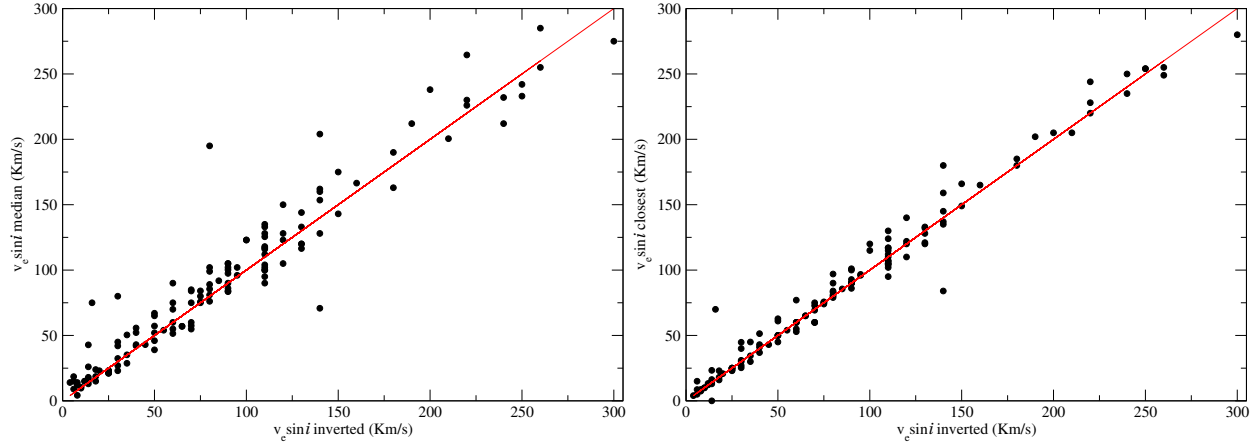


Fig. 6. *Left panel:* comparison between our inverted projected equatorial rotational velocities and the median of the catalogued ones. *Right panel:* comparison between our inverted projected equatorial rotational velocities and the closest of the catalogued $v \sin i$.

and 5 of Table 1). In comparison with the closest values, the absolute mean signed difference is found to be 1.7 km s^{-1} with a standard deviation of 10.0 km s^{-1} (Cols. 6 and 7 of Table 1). Taking into account the spread in the catalogued values by applying Eq. (11) decreases the average deviation to $\sim 2.0 \text{ km s}^{-1}$ and the corresponding standard deviation to 9.5 km s^{-1} (Cols. 2 and 3 of Table 1). This result shows that the projected equatorial rotational velocities of A-type stars can be derived with a small uncertainty using PCA-inversion techniques.

5.2.1. Well-studied A stars

We also checked the deviation of our derived parameters with those of some well-studied stars. To accomplish this, we selected all of the stars that have been studied extensively by different authors using different techniques (photometric, spectroscopic, and/or interferometric). In lack of a benchmark stars list for A stars (see for instance Blanco-Cuadros et al. 2014 for a list of FGK benchmark stars), our selection of well-studied stars is based on the number of references that we found and the number of derived values for each individual parameter for each star. In our selection, we also took stars with existing photometric determinations of $T_{\text{eff}}/\log g$ into account. We ended up with 19 stars with more than 120 references each. The selected stars are Vega, Sirius A, HD 22484, HD 15318, HD 76644, HD 49933, HD 214994, HD 214923, HD 113139, HD 114330, HD 27819, HD 5448, HD 33256, HD 29388, HD 91480, HD 30210, HD 32301, HD 28355, and HD 222603. Considering these stars, and comparing our inverted parameters to the median of the catalogued stars, we obtain an absolute mean signed difference of 35 K with a standard deviation of 250 K for T_{eff} , 0.18 dex with a standard deviation of 0.40 dex for $\log g$, 0.09 dex with a standard deviation of 0.15 dex for $[\text{Fe}/\text{H}]$, and 4.0 km s^{-1} with a standard deviation 9.5 km s^{-1} for $v \sin i$. These differences and standard deviations are shown in Cols. 8 and 9 of Table 1. The absolute mean signed differences and their corresponding standard deviations decrease when comparing with the closest catalogues values (Cols. 10 and 11 of Table 1). These calculations also show that the effective temperatures, metallicities, and projected equatorial rotational velocities of A stars can be derived with a good accuracy using PCA-inversion techniques.

6. Outliers

The outliers have been selected based on the large difference between the inverted parameters and the values retrieved from the

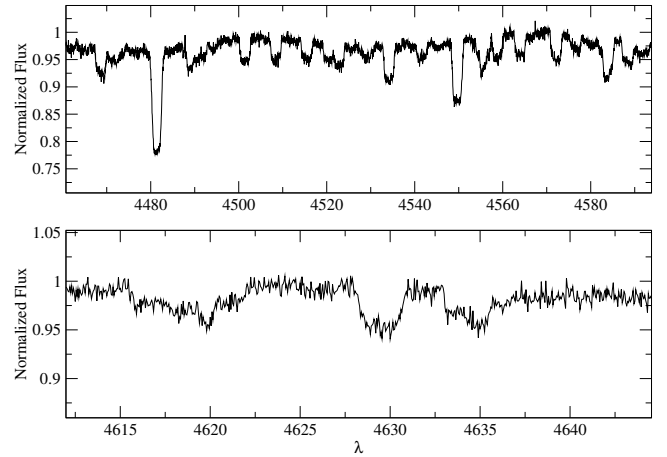


Fig. 7. Observed spectrum of HD 6530. Two different regions are shown in order to show the peculiar line profile that could be due to a binary system.

catalogues. We also considered as outliers, and checked in detail, stars that were found to have very low metallicities. As explained in Sect. 5, gravity is usually determined with poor accuracy in spectral fitting technique for $T_{\text{eff}} \leq 8000 \text{ K}$, but we are confident that with our large wavelength range we are able to reach realistic values for these stars due to the sensitivity of the metallic lines to surface gravity (Ryabchikova et al. 2015). This is justified by the small differences that we found between the inverted and median $\log g$ (Sect. 5.2). For that reason we decided to exclude gravity as a parameter for the outliers selection.

6.1. HD 6530

Cowley et al. (1969) classified this star as A1V with no observable peculiarity. Royer et al. (2014) suspected that this star could be a binary system as the cross-correlation function has an asymmetric profile. Figure 7 shows the profile of several lines in the SOPHIE spectrum of HD 6530. All lines have flat cores. More observations are needed to confirm whether HD 6530 is a binary system or whether signatures of gravity darkening due to fast rotation seen pole-on is present in the spectrum.

6.2. HD 12446

HD 12446 has been classified as a kA2hF2mF2 star by [Gray & Garrison \(1989\)](#) and is considered as a normal star in Simbad. [Eggleton & Tokovinin \(2008\)](#) found that this star is a member of a binary system containing a A0pSi primary ($V = 4.16$) and an A9 secondary ($V = 5.27$). The contribution of an Ap star and Am star to the spectrum probably accounts for the discrepancy between the inverted parameters and those retrieved from catalogues, especially for the projected equatorial rotational velocity.

6.3. HD 16605

HD 16605, member of NGC 1039, is an A1-type star according to [Renson & Manfroid \(2009\)](#). The magnetic field of the star was discovered by [Kudryavtsev et al. \(2006\)](#) with a longitudinal component that varies from -2430 G to -840 G. [Balega et al. \(2012\)](#) considered this star an A1p with a peculiarity of the type SiSrCr. They also spectroscopically derived an effective temperature of $10\,350$ K and a $v \sin i$ of 13 km s^{-1} close to our inverted values of $10\,800$ K and 14 km s^{-1} , respectively. Their values were not found by querying VizieR and, hence, are not present in Table 2. Using interferometric data, [Balega et al. \(2012\)](#) discovered that HD 16605 has a companion F star, 3.1 mag fainter, with an orbital period of 680 years. The fact that HD 16605 is an Ap star with overabundances of Si, Sr, and Cr could explain the large metallicity that we derived (1.1 dex). As an A1-type star, the catalogue median temperature of 8000 K is clearly underestimated. ATLAS9 model atmospheres and the spectrum synthesis used here do not properly account for the effect of the magnetic field. The magnetic pressure is ignored in the hydrostatic equilibrium when computing the atmospheric structure in ATLAS9 and the Zeeman splitting of the line profile is not included in the line synthesis. For that reason, the PCA-inversion technique presented in this work may not be applicable for Ap stars using ATLAS9 and SYNSPEC.

6.4. HD 23479

Classified as A7V, HD 23479 was suspected to be a spectroscopic binary by [Liu et al. \(1991\)](#). It has a visual companion in the Tycho Catalog with separation $0''.85$, $V = 9.45$, and $B - V = 0.54$ according to [Malkov et al. \(2012\)](#). These authors also showed that HD 23479 is a binary system containing a A7 + F6 star. Our inverted T_{eff} and $v \sin i$ (7300 K and 2.0 km s^{-1}) are in good agreement with previous findings (7239 K and 0.0 km s^{-1}). This star was flagged as an outlier because of the high iron abundance (+2.0 dex). This large overabundance could be explained by two contributors in the spectrum of this star, and as a spectroscopic binary, our inverted parameters do not represent the fundamental parameters of the components of HD 23479.

6.5. HD 50405

According to [McCuskey \(1956\)](#), HD 50405 is classified as a A0 star in Simbad. [Lefèvre et al. \(2009\)](#) detected, in their preliminary work, a possible variability in HD 50405 of the δ Scuti type. By carefully inspecting the NARVAL high-resolution spectrum of HD 50405 in Fig. 8, we found a contribution of another component. This star is a member of a binary system with a component of probably a similar spectral type and should be classified as such. The variation of the radial velocity between the two components was found to be constant along the spectrum with a

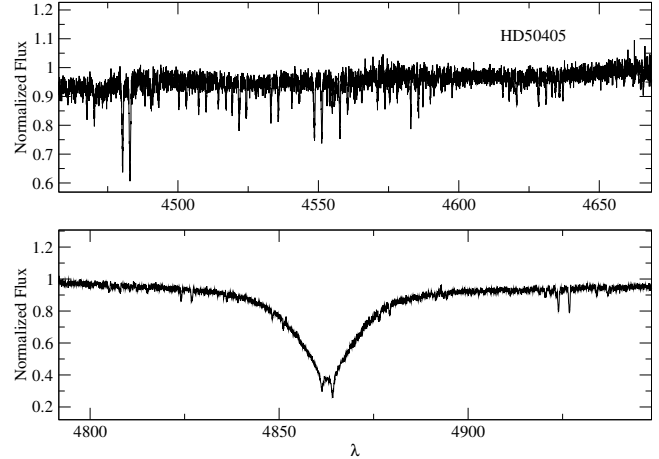


Fig. 8. Observed spectrum of HD 50405 around the Mg II and Fe II lines (upper panel) and the H β line (lower panel).

value close to 170 km s^{-1} . This strengthens the hypothesis that the star is actually a binary. This star will be subject of a detailed study to better characterize the system. The presence of double lines is the reason for the discrepancies between the inverted parameters and the catalogued parameters.

6.6. HD 172271

According to Simbad, this star is classified as A1 by [Renson & Manfroid \(2009\)](#). [Landstreet et al. \(2008\)](#) detected a magnetic field of ~ 260 G in this star and classified it as A0pCr. In their paper, they derived a $v \sin i$ of about 107 km s^{-1} very close to the 100 km s^{-1} that we found from our inversion. [Landstreet et al. \(2008\)](#) also derived the abundance of Cr and found an overabundance of about 2.5 dex, while Fe is about 0.5 dex overabundant. Their iron abundance is close to our inverted value of 0.4 dex. The inverted parameters of HD 172271 were rejected because it is an Ap star.

7. Discussion and conclusions

Using the PCA inversion, we have derived the effective temperature of “normal” A and Am stars with an average deviation of ~ 150 K with respect to the catalogued values and with a standard deviation of 500 K. We derive $\log g$, $[\text{Fe}/\text{H}]$ and $v \sin i$ with a deviation of 0.35 dex, 0.15 dex, and 2 km s^{-1} , respectively. Their standard deviations are 0.3 dex, 0.25 dex, and 9.5 km s^{-1} , respectively. This assessment was carried out using an automated and objective approach that takes the large spread in the catalogued data into account.

We also derived, for the first time, the metallicity of a large number of A-type stars. The combination of $v \sin i$ and $[\text{Fe}/\text{H}]$ is a very important parameter for a first detection of the Am phenomenon. It has been shown that for the same spectral type, Am stars have lower projected equatorial rotational velocities than “normal” A-type stars and are usually slightly overabundant in iron ([Gebran et al. 2008, 2010; Gebran & Monier 2008](#)).

One could also use a learning data set based on benchmark stars with accurate stellar parameters. This could be possible with large survey calibration such as *Gaia*. The radial velocity spectrometer (RVS) onboard of *Gaia* will collect about 15×10^6 spectra during its five-year missions. The RVS will provide spectra in the CaII IR triplet region (from 8470 to 8710 Å) at a spectral resolution of $\sim 11\,000$. In case of A-type stars, the

RVS wavelength range contains the CaII triplet as well as N I, S I, Si I, and the strong Paschen hydrogen lines that are a good indicator of surface gravity at the temperature of A-type stars. Recio-Blanco et al. (2016), using only 490 synthetic spectra as a learning database, showed that the parameters of A-type stars can be derived with high accuracy in the RVS spectral range (see Sect. 5.2 of their paper). In their work, they compared the performance of different codes on the stellar parameters derived from the RVS spectra. We intend to test the PCA-inversion technique in a similar manner, with a much larger database, to show its efficiency at low resolution and in the RVS spectral range for not only A-type stars, but also for late- and early-type stars.

A preliminary work was published in (Farah et al. 2015), where we estimated stellar parameters of early-type stars observed in the context of the *Gaia*-ESO Survey (GES; Gilmore et al. 2012). This survey collects observations of faint stars ($14 < V < 19$) taken using the GIRAFFE/FLAMES spectrograph at a medium resolution of $R \sim 25\,000$. Two spectral regions were used, one covering the H δ line and one between 4400 and 4550 Å, that are similar to the wavelength range used in the present work. Although the spectral resolution is low compared to those of the SOPHIE or PolarBase spectra, we derived stellar parameters in good agreement with those recommended by the GES community.

As a conclusion, the work of Paletou et al. (2015b,c) and the present one show that the PCA-inversion method proves to be fast and efficient for inverting stellar parameters of FGK and A/Am stars and deriving effective temperatures of M dwarf stars. We built a learning database containing 2.2×10^6 synthetic spectra at each resolution. For the first time, we used the power iteration algorithm to derive the eigenvectors/principal components of such a large database. The advantage of such an algorithm is that it can be used no matter the size of the database. All the synthetic data were calculated with a microturbulence velocity of 2 km s^{-1} . This parameter, in case of 1D models, should be considered with caution, especially when attempting to model line profiles. One should modify this parameter in the learning database adopting the appropriate value for the effective temperature of the star. Gebran et al. (2014) derived the dependence of ξ_t on the effective temperature showing a broad maximum around 8000 K. We are working on implementing this equation in the calculations of learning databases for F-A-B stars.

Acknowledgements. This research has made use of the VizieR catalogue access tool, CDS, Strasbourg, France. The original description of the VizieR service was published in A&AS 143, 23. This research has made use of the SIMBAD database, operated at CDS, Strasbourg, France. This research has used the Polarbase, Elodie, and Sophie databases. PolarBase is operated by the OV-GSO data centre, CNRS-INSU, and Observatoire Midi-Pyrénées – Université Paul Sabatier, Toulouse, France (Paletou et al. 2015a). M.G. thanks Dr. Joe Malkoun for his valuable comments and help in the implementation of the power iteration technique. We also thank the anonymous referee for his valuable comments that helped improving this manuscript.

References

- Bailer-Jones, C. A. L., Irwin, M., & von Hippel, T. 1998, *MNRAS*, **298**, 361
 Balega, Y. Y., Dyachenko, V. V., Maksimov, A. F., et al. 2012, *Astrophys. Bull.*, **67**, 44
 Baranne, A., Queloz, D., Mayor, M., et al. 1996, *A&AS*, **119**, 373
 Blanco-Cuadros, S., Soubiran, C., Jofré, P., & Heiter, U. 2014, *A&A*, **566**, A98
 Borne, K. 2013, *Planets, Stars and Stellar Systems*, Vol. 2: Astronomical Techniques, Software and Data, 403
 Castelli, F., & Kurucz, R. L. 2003, in *Modelling of Stellar Atmospheres*, *Proc. IAU Symp.*, **210**, 20
 Catanzaro, G., & Balona, L. A. 2012, *MNRAS*, **421**, 1222
 Code, A. D., Bless, R. C., Davis, J., & Brown, R. H. 1976, *ApJ*, **203**, 417
 Cowley, A., Cowley, C., Jaschek, M., & Jaschek, C. 1969, *AJ*, **74**, 375
 Demmel, J. W. 1997, *Applied Numerical Linear Algebra* (Philadelphia: SIAM)
 Eggleton, P. P., & Tokovinin, A. A. 2008, *MNRAS*, **389**, 869
 Farah, W., Gebran, M., Paletou, F., & Blomme, R. 2015, *Proc. SF2A*, **3**, <http://sf2a.eu/spip/spip.php?article638>
 Gazzano, J.-C., de Laverny, P., Deleuil, M., et al. 2010, *A&A*, **523**, A91
 Gebran, M., & Monier, R. 2008, *A&A*, **483**, 567
 Gebran, M., Monier, R., & Richard, O. 2008, *A&A*, **479**, 189
 Gebran, M., Vick, M., Monier, R., & Fossati, L. 2010, *A&A*, **523**, A71
 Gebran, M., Monier, R., Royer, F., Lobel, A., & Blomme, R. 2014, *Putting A Stars into Context: Evolution, Environment, and Related Stars*, 193
 Gilmore, G., Randich, S., Asplund, M., et al. 2012, *The Messenger*, **147**, 25
 Gray, D. F. 1992, *The observation and analysis of stellar photospheres*, *Camb. Astrophys. Ser.*, **20**
 Gray, R. O., & Garrison, R. F. 1989, *ApJS*, **70**, 623
 Grevesse, N., & Sauval, A. J. 1998, in *Solar Composition and Its Evolution – From Core to Corona*, *Proc. ISSI Workshop*, 161
 Heiter, U., Kupka, F., van't Veer-Menneret, C., et al. 2002, *A&A*, **392**, 619
 Hill, G. M., & Landstreet, J. D. 1993, *A&A*, **276**, 142
 Hill, G., Gulliver, A. F., & Adelman, S. J. 2010, *ApJ*, **712**, 250
 Hubeny, I., & Lanz, T. 1992, *A&A*, **262**, 501
 Jolliffe, I. T. 1986, *Springer Series in Statistics* (Berlin: Springer)
 Kantor, J., Axelrod, T., Becla, J., et al. 2007, *Astronomical Data Analysis Software and Systems XVI*, **376**, 3
 Kudryavtsev, D. O., Romanyuk, I. I., Elkin, V. G., & Paunzen, E. 2006, *MNRAS*, **372**, 1804
 Kurucz, R. L. 1992, *Rev. Mex. Astron. Astrofis.*, **23**, 45
 Kurucz, R. L. 2005, *Mem. Soc. Astron. It. Suppl.*, **8**, 14
 Landstreet, J. D. 2011, *A&A*, **528**, A132
 Landstreet, J. D., Silaj, J., Andretta, V., et al. 2008, *A&A*, **481**, 465
 Lefèvre, L., Michel, E., Aerts, C., et al. 2009, *Comm. Asteroseismol.*, **158**, 189
 Liu, T., Janes, K. A., & Bania, T. M. 1991, *ApJ*, **377**, 141
 Malkov, O. Y., Tamazian, V. S., Docobo, J. A., & Chulkov, D. A. 2012, *A&A*, **546**, A69
 McCuskey, S. W. 1956, *ApJS*, **2**, 271
 Moulata, J., Ilvaysky, S. A., Prugniel, P., & Soubiran, C. 2004, *PASP*, **116**, 693
 Munari, U., Zwitter, T., & Siebert, A. 2005, *The Three-Dimensional Universe with Gaia*, *ESA SP*, **576**, 529
 Paletou, F., & Zolotukhin, I. 2014, *ArXiv e-prints* [[arXiv:1408.7026](https://arxiv.org/abs/1408.7026)]
 Paletou, F., Glorian, J.-M., Génot, V., et al. 2015a, in *SF2A-2015: Proc. Annual meeting of the French Society of Astronomy and Astrophysics*, eds. F. Martins, S. Boissier, V. Buat, L. Cambrézy, & P. Petit, 37
 Paletou, F., Böhm, T., Watson, V., & Trouillet, J.-F. 2015b, *A&A*, **573**, A67
 Paletou, F., Gebran, M., Houdebine, E. R., & Watson, V. 2015c, *A&A*, **580**, A78
 Perryman, M. A. C., de Boer, K. S., Gilmore, G., et al. 2001, *A&A*, **369**, 339
 Petit, P., Louge, T., Théado, S., et al. 2014, *PASP*, **126**, 469
 Przybilla, N., & Butler, K. 2004, *ApJ*, **609**, 1181
 Recio-Blanco, A., de Laverny, P., Allende Prieto, C., et al. 2016, *A&A*, **585**, A93
 Rees, D. E., López Ariste, A., Thatcher, J., & Semel, M. 2000, *A&A*, **355**, 759
 Re Fiorentin, P., Bailer-Jones, C. A. L., Lee, Y. S., et al. 2007, *A&A*, **467**, 1373
 Renson, P., & Manfroid, J. 2009, *A&A*, **498**, 961
 Richer, J., Michaud, G., & Turcotte, S. 2000, *ApJ*, **529**, 338
 Richard, O., Michaud, G., & Richer, J. 2001, *ApJ*, **558**, 377
 Richard, O., Michaud, G., & Richer, J. 2002, *ASP Conf. Proc.*, **185**: *Radial and Nonradial Pulsations as Probes of Stellar Physics*, 259, 270
 Roweis, S. 1998, in *Advances in Neural Information Processing Systems* (MIT Press), 626
 Royer, F., Grenier, S., Baylac, M.-O., Gómez, A. E., & Zorec, J. 2002, *A&A*, **393**, 897
 Royer, F., Gebran, M., Monier, R., et al. 2014, *A&A*, **562**, A84
 Ryabchikova, T., Piskunov, N., & Shulyak, D. 2015, *Physics and Evolution of Magnetic and Related Stars*, **494**, 308
 Smalley, B. 2004, *The A-Star Puzzle*, *IAU Symp.*, **224**, 131
 Smalley, B. 2005, *Mem. Soc. Astron. It. Suppl.*, **8**, 130
 Smalley, B., & Dworetzky, M. M. 1993, *A&A*, **271**, 515
 Steinmetz, M. 2003, *GAIA Spectroscopy: Science and Technology*, *ASP Conf. Proc.*, **298**, 381
 Takeda, Y., Kawanomoto, S., & Ohishi, N. 2008, *ApJ*, **678**, 446
 Takeda, Y., Kang, D.-I., Han, I., Lee, B.-C., & Kim, K.-M. 2009, *PASJ*, **61**, 1165
 Tonry, J., & Davis, M. 1979, *AJ*, **84**, 1511
 Vick, M., Michaud, G., Richer, J., & Richard, O. 2010, *A&A*, **521**, A62
 York, D. G., Adelman, J., Anderson, J. E., Jr., et al. 2000, *AJ*, **120**, 1579
 Zahn, J.-P. 2005, *EAS Pub. Ser.*, **17**, 157

Sliced Inverse Regression : application to stellar fundamental parameters

Résumé : Nous présentons une nouvelle méthode pour la détermination des paramètres stellaires fondamentaux. La procédure est basée sur une régression inverse par tranches régularisée (RSIR). Nous avons d’abord testé l’efficacité de cette procédure sur un échantillon synthétique bruité d’étoiles A, F, G et K, et avons inversé simultanément les paramètres fondamentaux : température effective T_{eff} , gravité de surface $\log(g)$, métallicité $[M/H]$ et en plus, la vitesse de rotation projetée $v \sin(i)$. Différentes bases de données d’apprentissage ont été calculées pour les différents types d’étoiles, et plusieurs tests ont été effectués en utilisant des bases de données avec différents pas d’échantillonnage pour T_{eff} , $\log(g)$, $v \sin(i)$ et $[M/H]$. Nous montrons que la troncature des composantes principales permet, en sélectionnant un ensemble de plus proches voisins, la constitution d’une nouvelle base d’apprentissage pour l’inversion avec RSIR. Pour tous les types spectraux, diminuer la taille de la base d’apprentissage nous permet d’atteindre de meilleures précisions internes qu’en utilisant l’analyse en composantes principales sur de plus grandes bases d’apprentissage. Pour chaque paramètre analysé nous avons atteint des erreurs internes plus faibles que le pas d’échantillonnage du paramètre. Nous avons aussi appliqué la technique sur des étoiles de types FGK et A observées. Les paramètres inversés des étoiles observées sont en accord avec ceux obtenus dans les études antérieures, mais avec de meilleures précisions. La technique d’inversion RSIR, complétée par un pré-traitement basé sur l’analyse en composantes principales, se montre comme étant un outil rapide et efficace pour l’estimation des paramètres stellaires des étoiles de types A, F, G et K.

(Soumis le 1/05/2018 à MNRAS)

Research Article

S. Kassounian¹, M. Gebran^{1*}, F. Paletou², and V. Watson²

Sliced Inverse Regression: application to stellar fundamental parameters

DOI: DOI

Received ...; revised ...; accepted ...

Abstract: We present a method for the determination of stellar fundamental parameters. The procedure is based on a regularized sliced inverse regression (RSIR). We first tested the effectiveness of this procedure on a noisy synthetic sample of A, F, G, and K stars, and inverted simultaneously the fundamental parameters: effective temperature T_{eff} , surface gravity $\log g$, metallicity $[M/H]$ and, in addition, the projected rotational velocity $v \sin i$. Different learning databases were calculated for different types of stars, and several tests were done using databases with different sampling in T_{eff} , $\log g$, $v \sin i$, and $[M/H]$. The use of principal component analysis (PCA) allows in decreasing the dimension of the spectra. Combined with the nearest neighbor (NN) search, the PCA reduced the size of the learning database by selecting the useful spectra for the inversion. Tikhonov regularization helps in dealing with the ill-conditioning problem existing in SIR. For all spectral types, decreasing the size of the learning database allowed us to reach internal accuracies better than the one reached using only PCA NN search with larger learning databases. For each analyzed parameter, we have reached internal errors that are smaller than the sampling step of the parameter. We have also applied the technique to a sample of observed FGK and A stars. For a selection of well studied stars, the inverted parameters are in agreement with the ones derived in previous studies. The RSIR inversion technique, complemented with a PCA pre-processing steps proves to be an efficient tool for estimating stellar parameters of A, F, G, and K stars.

Keywords: methods: data analysis, methods: statistical, techniques: spectroscopic, stars: fundamental parameters.

1 Introduction

Astronomical surveys, either spaceborne or ground-based, are gathering an unprecedented amount of data. One can mention the SDSS DR14 data (Abolfathi et al., 2017) that contains 154TB of millions of spectroscopic and photometric data. The DR5 of the LAMOST survey (Cui et al., 2012) contains 9 million spectra in total. Gaia DR2 (Perryman et al., 2001) provides information about 1.3 billion stars (Katz & Brown, 2017). These space and on-ground surveys quantify the size of the data the astronomical community will face in a near future.

Spectroscopic analysis is crucial for the derivation of stellar fundamental atmospheric parameters which are the effective temperature (T_{eff}), the surface gravity ($\log g$), and the metallicity ($[M/H]$). In addition to these fundamentals, and because it may strongly affect the shape of the observed spectra, the projected equatorial rotational velocity, $v \sin i$, is also retrieved from spectroscopic information. Many authors have for long been using spectroscopic data to estimate the stellar atmospheric

parameters (Buchhave et al., 2012, Dieterich et al., 2017, Fabbro et al., 2018, Latham et al., 2002, McWilliam, 1990, Schönrich & Bergemann, 2014, Torres et al., 2002). In order to extract the most relevant and accurate information from high-resolution stellar spectra, still more endeavour is required.

Most of the traditional approaches and developed pipelines rely on standard procedures such as comparing an observed spectrum with a set of theoretical spectra (Morris et al., 2018, Valenti & Piskunov, 1996). The requirement for advanced computational techniques rises from the generated large dimensionality of the data due to the wide wavelength coverage together with high spectral resolution. Many new techniques are being developed based on statistics, dimensionality reduction, or processing of observational data. In Ness et al. (2015) and Casey et al. (2016), a data-driven approach is introduced (CANNON) for determining stellar labels (fundamental parameters and detailed stellar abundances) from spectroscopic data. Their learning databases (LDB) are based on a subset of reference objects for which the stellar labels are known with high accuracy. Dimension reduction tech-

niques are also developed and used, such as applying the Principal Component Analysis (PCA) for data reduction (see e.g. Jolliffe 1986). PCA has shown its effectiveness in inverting the stellar fundamental atmospheric parameters in several studies (Bailer-Jones et al., 1998, Gebran et al., 2016, Paletou et al., 2015a,b, Re Fiorentin et al., 2007). Xiang et al. (2017) estimated the stellar atmospheric parameters as well as the absolute magnitudes and α -elements abundances from the LAMOST spectra with a multivariate regression method based on kernel-based PCA. The LAMOST spectroscopic survey data has also been recently analyzed by Boeche et al. (2018) to invert stellar parameters and chemical abundances in which several combined approaches and techniques were compared. The authors developed a code called SP_Ace which utilizes nearest neighbor comparison and non-linear model fitting techniques. In Wilkinson et al. (2017), a spectral fitting code (FIREFLY) was developed to derive the stellar population properties of stellar systems. FIREFLY uses a χ -squared minimization fitting procedure that fits stellar population models to spectroscopic data, following an iterative best-fitting process controlled by a Bayesian information criterion. Their approach is efficient to overcome the so-called “ambiguities” in the spectra. More recently, Gill et al. (2018) used wavelet decomposition to distinguish between noise, continuum trends, and stellar spectral features in the CORALIE FGK-type spectra. By calculating a subset of wavelet coefficients from the target spectrum and comparing it to those from a grid of models in a Bayesian framework, they were able to derive T_{eff} , $[M/H]$, and $v \sin i$ for these stars.

In this study, we apply a technique, that uses the PCA reduction and nearest neighbor search (Gebran et al., 2016, Paletou et al., 2015a,b) complemented with a Regularized Sliced Inverse Regression (Bernard-Michel et al., 2007, 2009) (RSIR) procedure in order to derive simultaneously T_{eff} , $\log g$, $[M/H]$ and $v \sin i$ from spectra of A, and FGK type stars. Up to now, sliced inverse regression has been rarely used in astronomy (Bernard-Michel et al., 2009, Watson et al., 2017). When combined with PCA, the derivations of the fundamental atmospheric parameters are achieved with higher accuracy compared to the PCA nearest neighbor inversion (Gebran et al., 2016, Paletou et al., 2015a,b) alone. The mathematical description of our method is detailed in Sec. 2. Section 3 describes the elements used for the enhancement of the computational abilities of SIR. Section 4 discusses the application of the technique on stellar synthetic data as well as the results for A, F, G, and K simulated stars. In

section 5, we show the results of inversions of real stars. Discussion and conclusion are gathered in Sec. 6.

2 Sliced inverse regression (SIR)

SIR, originally formulated by Li (1991), is a statistical technique that reduces multivariate regression to a lower dimensional. It finds an inverse functional relationship between the response and the predictor which are the fundamental parameters and the flux respectively. Synthetic spectra flux values, x_{syn} , are usually calculated based on the set of stellar atmospheric parameters in the form of:

$$x_{syn} = f(T_{\text{eff}}, \log g, [M/H], v \sin i). \quad (1)$$

The inverse functional relation is used to predict the parameters of the observed flux values, x_{obs} , in the form of:

$$f^{-1}(x_{obs}) = (T_{\text{eff}}, \log g, [M/H], v \sin i). \quad (2)$$

In our work, we have derived a functional relationship for each parameter in the following way:

$$Y_j = f_j^{-1}(x_{obs}), \quad (3)$$

where $j = 1, 2, 3, 4$ for T_{eff} , $\log g$, $[M/H]$, and $v \sin i$.

2.1 Global covariance matrix Σ

SIR starts from the computation of the covariance matrix Σ of all the synthetic spectra x_i of the LDB. The spectra are represented in a matrix of dimension $N_\lambda \times N_{\text{spectra}}$, where N_λ is the number of wavelength points per spectrum and N_{spectra} is the total number of spectra in the LDB. The covariance matrix Σ , is in the following way:

$$\Sigma = \frac{1}{N_{\text{spectra}}} \sum_{i=1}^{N_{\text{spectra}}} (x_i - \bar{x}) \cdot (x_i - \bar{x})^T, \quad (4)$$

where the global mean $\bar{x}$ is defined as:

$$\bar{x} = \frac{1}{N_{\text{spectra}}} \sum_{i=1}^{N_{\text{spectra}}} x_i, \quad (5)$$

x_i being a row vector containing the flux values of spectrum i .

2.2 Intra-slices covariance matrix

The following step in the SIR procedure is to organize the entire spectra and stack them in an increasing order

of the considered parameters for inversion. For example, if we are to invert T_{eff} of each star, the spectra database should be organized in increasing order of T_{eff} having the other parameters ordered randomly.

We then selected sets of spectra, also called “slices”, having the same considered stellar parameters. These slices do not overlap each other (Li, 1991). Then we calculate the means $\bar{x}_h$ of the slice of the spectra found in each slice S_h that contains n_h synthetic spectra (h being the index of each slice). For the inversion of each parameter, $\bar{x}_h$ and $\bar{x}$ are used to calculate the “intra-slices” covariance matrix, Γ :

$$\Gamma = \sum_{h=1}^H \frac{n_h}{N} (\bar{x}_h - \bar{x}) \cdot (\bar{x}_h - \bar{x})^T, \quad (6)$$

where

$$\bar{x}_h = \frac{1}{n_h} \sum_{x \in S_h} x_i. \quad (7)$$

2.3 Dimension reduction and parameter inversion

SIR builds a subspace that maximizes the variance between the slices while minimizing the variance within the slices. This subspace insures a higher accuracy of regressive prediction by bringing together elements which have close parameter values (Watson et al., 2017). On the other hand, slicing the spectra based on non-overlapping similar parameters insures this inter-slice maximization and intra-slice minimization.

Dimension reduction starts by calculating the matrix $\Sigma^{-1}\Gamma$ where Σ and Γ are the two previously defined matrices. One eigenvector of $\Sigma^{-1}\Gamma$, called β_λ and corresponding to an eigenvalue λ , is used to form the reduction subspace. This will allow us to do regression in a 2-dimensional space using an inverse functional relationship. This relationship is constructed via a linear piecewise interpolation between the projection coordinates of the slices on the single eigenvector of $\Sigma^{-1}\Gamma$ and the parameters.

The selection of β_λ , is based on a metric C_λ that quantifies the relationship between the spectra and the parameters. C_λ , defined as the “sliced inverse regression criteria” (Bernard-Michel et al., 2007, 2009), is calculated as follows:

$$C_\lambda = \frac{\beta_\lambda^t \Gamma \beta_\lambda}{\beta_\lambda^t \Sigma \beta_\lambda} \approx \frac{\text{Var}(\langle \beta_\lambda \cdot x_i \rangle)}{\text{Var}(\langle \beta_\lambda \cdot x_i \rangle) + \text{Var}(S_h)}, \quad (8)$$

where β_λ^t is the transpose of β_λ and Var is the classical variance function.

$\beta_\lambda^t \Gamma \beta_\lambda$ is considered to be as the “inter-slice” variance, whereas $\beta_\lambda^t \Sigma \beta_\lambda$ represents the total variance which is interpreted as the sum of the “inter-slice” variance and the “intra-slice” variance. The β_λ that gives a C_λ value closest to 1 is considered as the best choice for the reducing basis vector. In the present work, C_λ varies between 0.91 and 0.97 when using the eigenvector of $\Sigma^{-1}\Gamma$ with the largest eigenvalue λ .

To invert the parameters, we apply linear piecewise interpolation on the coordinates of the projections of the $\bar{x}_h$ -s on β_λ . The estimation of the parameters is derived according to:

$$\hat{y} = \begin{cases} \bar{y}_1, & \text{if } x^p \in]-\infty, \bar{x}_1^p], \\ \bar{y}_h + \left(x_{\text{obs}}^p - \bar{x}_h^p \right) \left(\frac{\bar{y}_{h+1} - \bar{y}_h}{\bar{x}_{h+1}^p - \bar{x}_h^p} \right), & \text{if } x^p \in]\bar{x}_h^p, \bar{x}_{h+1}^p], \\ \bar{y}_H, & \text{if } x^p \in]\bar{x}_H^p, +\infty[, \end{cases} \quad (9)$$

where $\hat{y}$ is the estimated parameter; $\bar{y}_h$ is the mean of the parameters of the spectra in slice h . The superscript “ p ” represents the projected value of a selected set of data on β_λ ($x^p = \langle \beta_\lambda \cdot x \rangle$).

3 Enhancement of the computational abilities of SIR

In the present work, we are dealing with large amounts of high resolution spectra, where $\Sigma^{-1}\Gamma$ has a dimension of $\sim 10^4 \times 10^4$. In addition, using a large LDB for SIR induces an increase in the intra-slice variance. This will lead to less accurate inverted parameters. Therefore to simultaneously address these problems, we applied 2 additional steps to SIR: first, using PCA, we reduced the dimension of the synthetic and observed spectra (Watson et al., 2017) from $\sim 10^4$ to 12. Second, we applied NN search in the reduced subspace to select a specific LDB, which is smaller and represents the observed spectrum in a better way.

The ill-conditioning of $\Sigma^{-1}\Gamma$ can be justified by calculating the condition number. The higher the value of this number is, the more noise sensitive the system becomes (Kreyszig, 2010). For the present work, its value has reached as high as 10^{20} . In that case β_λ has found to be very noise sensitive leading to an unstable functional relationship. As a result, inaccurate inverted parameters are derived. To solve this issue, we have applied Tikhonov

regularization which aims to study the noise of each observed spectrum and add a priori information to $\Sigma^{-1}\Gamma$. Finally, the final scheme of all these steps are displayed in Fig. 1 and discussed in secs. 3.1 to 3.3.

3.1 Dimension reduction via PCA

PCA is a numerical technique that reduces the dimension of the data (spectra in our case) by projecting it on a set of orthogonal basis vectors called *principal components* or PC's. These components are the eigenvectors of the global covariance matrix Σ . The number of required PC's is determined based on an error reconstruction minimization. Paletou et al. (2015a) and Gebran et al. (2016) showed that for their databases only 12 PC's associated to the largest eigenvalues are enough to reduce the LDB. The reconstruction error for that case was less than 1%. For our tests, we have used the same LDB's (i.e. same range of wavelengths, fundamental stellar parameters, and resolution) as that of Paletou et al. (2015a) and Gebran et al. (2016). Therefore the use of 12 PC's is mathematically justified. As a result, the new LDB has a dimension of $N_{spectra} \times 12$ (See Fig. 1).

3.2 Tikhonov regularization

Tikhonov regularization (Vogel, 2002) is a technique that deals with ill-conditioned systems. It inserts a regularization parameter $\delta > 0$ to that system based on a priori information. As explained previously, $\Sigma^{-1}\Gamma$ is ill-conditioned, therefore the regularization process starts by calculating:

$$(\Sigma^2 + \delta I)^{-1} \Sigma \Gamma. \quad (10)$$

The eigenvector $\beta_\lambda(\delta)$ of the largest eigenvalue of the matrix defined in Eq. 10 is calculated based on an optimization approach. For each parameter of the observed star, an optimum δ is derived based on an intrinsic noise information from its spectrum. This procedure is initiated by estimating the signal to noise ratio (S/N) of the observed spectrum using the procedure of Stoehr et al. (2008). Then a random set of synthetic spectra are selected from the new LDB, and Gaussian white noise is added to them. SIR is finally applied to this selected random set and the prediction of their parameters is done via the piecewise interpolation process described in Eq. 9.

This internal inversion for different δ values leads to the selection of an optimum $\beta_\lambda(\delta)$.

δ is estimated by minimizing the difference between the newly inverted parameter values of the randomly selected noise added spectra ($\hat{y}_i$) and their initial noiseless values (y_i). The comparison is done using a normalized χ^2 :

$$\chi_N^2 = \sqrt{\frac{\sum_{i=1}^n (\hat{y}_i - y_i)^2}{\sum_{i=1}^n (y_i - \bar{y}_i)^2}}. \quad (11)$$

It was found that $\log_{10}(\chi_N^2)$ as a function of $\log(\delta)$ is a unimodal function which has a local minimum. This function is displayed in Figure 3 for a synthetic spectra having T_{eff} , $\log g$, $[M/H]$, $v \sin i$ and S/N of 7600 K, 2.50 dex, 0.0 dex, and 197 km s^{-1} , and 196, respectively. The original LDB used in this example is the one of Gebran et al. (2016), explained in detail in Sec. 4. To find the minimum of these curves, we applied a golden-section search algorithm (Kiefer, 1953). It is a classical numerical technique that minimizes unimodal functions which have a global minimum. The inversion process for each analyzed spectrum, and each parameter, has its own $\chi_N^2 = f(\delta)$ that needs to be minimized. From now on this approach will be known as regularized sliced inverse regression (RSIR).

3.3 Integrated scheme of the enhancements

Now that we have described the tools that were used to improve the procedure of SIR, we discuss how these techniques are integrated to increase the accuracy of the inversion process. The flowchart in Fig. 1 summarizes our adopted approach. The original LDB may reach to a dimension of $N_{spectra} \times N_\lambda \simeq 10^6 \times 10^4$. This is due to the fine sampling in the parameters, the high dimension of the spectra, and the large wavelength range which makes the process of SIR computationally heavy in terms of memory and time. In that case Σ and $\Sigma^{-1}\Gamma$ could reach to a dimension of order $\sim 10^6 \times 10^4$.

The first enhancement of the SIR is applying PCA to reduce the dimension of the LDB. This will transform the original LDB to a newly reduced one of dimension $N_{spectra} \times 12$ (see Fig 1). Second, using this newly reduced LDB composed of coefficients, a NN search was performed for each observed spectrum using a least square distance $d_j^{(O)}$, defined as:

$$d_j^{(O)} = \sqrt{\sum_{k=1}^{12} (\varrho_k - p_{jk})^2}, \quad (12)$$

where ϱ_k is the projection coordinate on the k^{th} dimension for an observed spectrum, and p_{jk} is the projection coefficient on the k^{th} dimension for the j^{th} synthetic spectrum.

During the inversion at least 2 distinct parameter values for each slice are required to construct the functional relationship. Therefore to select the optimum reduced LDB, a test for the construction of this relationship is required. Iteratively we tested for the number of distinct parameters by increasing the number of spectra of the nearest neighbors. Whenever the values of the distinct concerned parameters become greater or equal to 2, the inversion proceeds. This iterative approach also removes the case of singularity of $[\Sigma^2 + \delta I]^{-1}$ by adding additional nearest neighbors. Using optimum reduced LDB's which contain spectra that are closer to the observed ones theoretically insures a higher accuracy of inversion in SIR compared to using the entire original LDB. This enhancement is linked to the core concept of the SIR, where by selecting a set of nearest neighbors, we insure a lower minimization value of the intra-slice variance $Var(S_h)$ in Eq. 8. Now within each slice the spectra are closer to each other and they are closer to the average spectrum of the slice. At the same time, choosing these optima reduced LDB's overcomes the existence of any degeneracies by merging them into the slices.

Third, $\Sigma^{-1}\Gamma$ now has a dimension of 12×12 and is constructed from the optimized reduced LDB. Its high condition number implies that it is ill-conditioned (see the example of the left panel of Fig. 2). Therefore to improve the inversion process for each observed spectrum, we apply the Tikhonov regularization in SIR for our selected optima reduced LDB's, iteratively. By applying this regularization, we are effectively taking advantage of the S/N ratio analysis and inserting the propagated noise information as a priori. In other words, we are applying an internal de-noising procedure.

In Fig. 2, we display the inversion results for the T_{eff} of a noisy synthetic spectrum. This spectrum has a T_{eff} value of 7600 K with an added Gaussian white noise of $S/N = 196$. As we iterated over different sizes of optimized reduced LDB's, a stability can be noted for the inverted values. In all our tests, we noticed that the number of spectra in the optimized reduced LDB's did not surpass 500. It is shown in this figure that the condition number of the non-regularized matrix is ~ 5 orders of magnitude larger than the ones in which the Tikhonov regularization was applied. The right panel displays the effect of the regularization on the inverted parameter (T_{eff}) of

the same spectrum. It is clearly shown that whatever the number the nearest neighbors in the optimized reduced LDB is, inversion is achieved with higher accuracy than the one without regularization.

Figure 3 represents the minimization of the $\log_{10}(\chi_N^2)$ as a function of $\log_{10}(\delta)$ for different sets of optimized reduced LDB's. This figure shows the unimodal nature of the curves irrespective of the size of the LDB.

4 Simulations and tests

In this section, we present the implementation and results of RSIR for two different sets of synthetic. We also compare these results to the ones of the PCA NN to show the improvement in the accuracies in the derived parameters. To each of these spectra, white Gaussian noise was added with a random S/N. The spectra were calculated in the range of A to K type stars. The reason for selecting this spectral range is that in Sec. 5, we apply this procedure to a sample of the observed stars studied in Paletou et al. (2015a) and Gebran et al. (2016).

4.1 The learning databases

As done in Paletou et al. (2015b) and Gebran et al. (2016), we have used model atmospheres were calculated using ATLAS9 with the new opacity distribution function (Castelli & Kurucz, 2003, Kurucz, 1992). These models assume local thermodynamic equilibrium (LTE), hydrostatic equilibrium, and a 1D plane-parallel atmosphere. Convection was treated using a mixing length parameter of 0.5 for $7000 \text{ K} \leq T_{\text{eff}} \leq 8500 \text{ K}$, and 1.25 for $T_{\text{eff}} \leq 7000 \text{ K}$, following the prescriptions of Smalley (2004). Synthetic spectra were calculated using SYNSPEC48 (Hubeny & Lanz, 1992). The adopted line lists were from Kurucz ghyperall.dat¹ and modified with more recent and accurate atomic data retrieved from the VALD² and the NIST³ databases (for more details see Gebran et al. 2016).

¹ <http://kurucz.harvard.edu>

² <http://www.astro.uu.se/~vald/php/vald.php>

³ <http://physics.nist.gov>

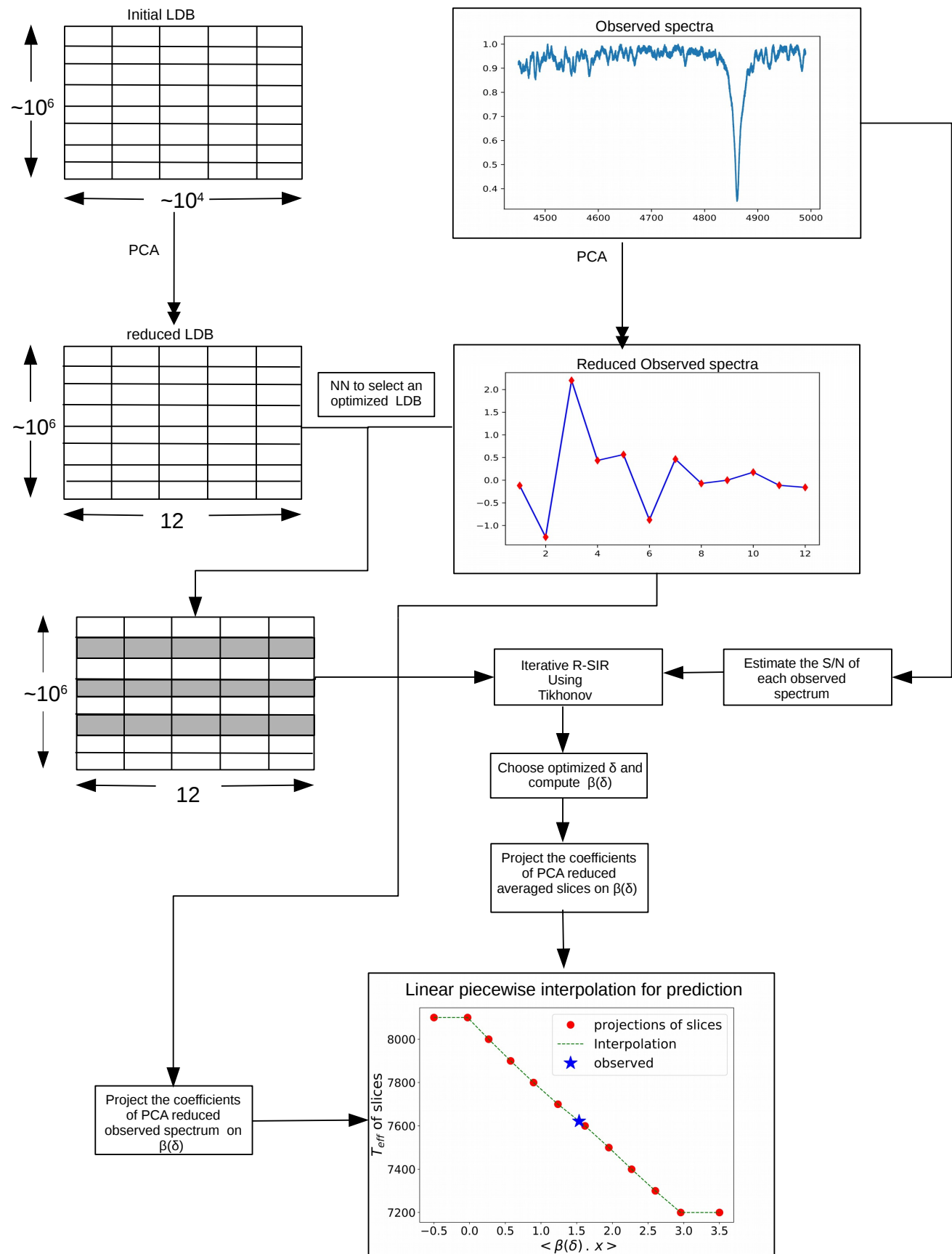


Fig. 1. Flowchart of the procedure. The numerics in this figure are for the inversions of the test described in Sec. 3.

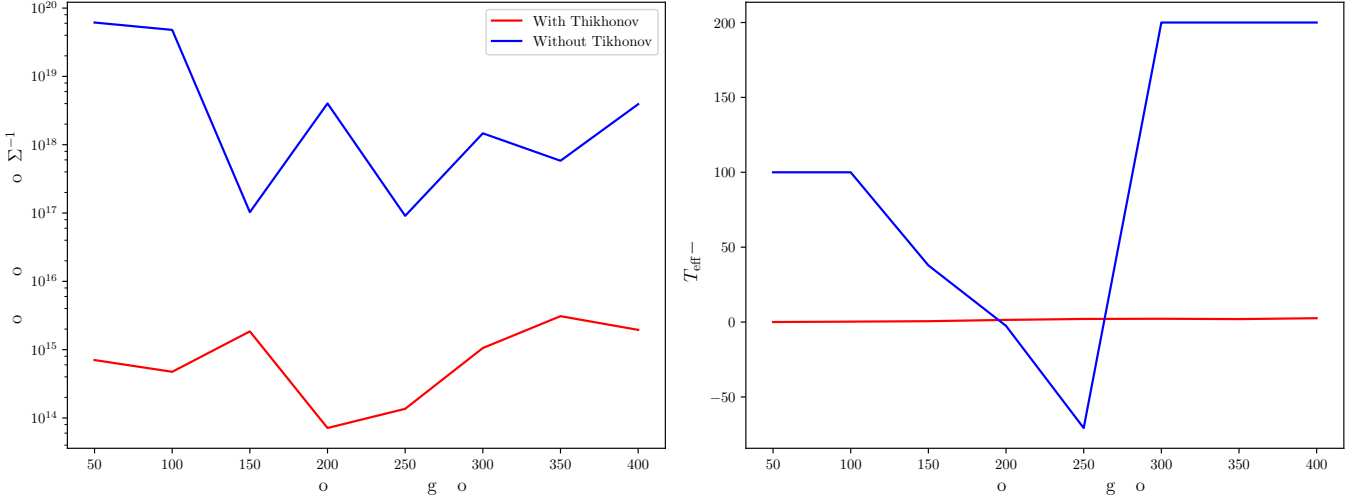


Fig. 2. Left: variation of the condition number as a function of the number of nearest neighbor using the optimum reduced LDB with and without the Tikhonov regularization. Right: the inverted T_{eff} as a function of the number of nearest neighbor with and without Tikhonov regularization. The inversion is done for a noise added synthetic spectrum with $T_{\text{eff}}=7600$ K, $\log g=2.50$ dex, $[M/H]=0$ dex, $v \sin i=197 \text{ km s}^{-1}$ and $S/N=196$. For clarity we display on the value of $T_{\text{eff}}-7600$ K

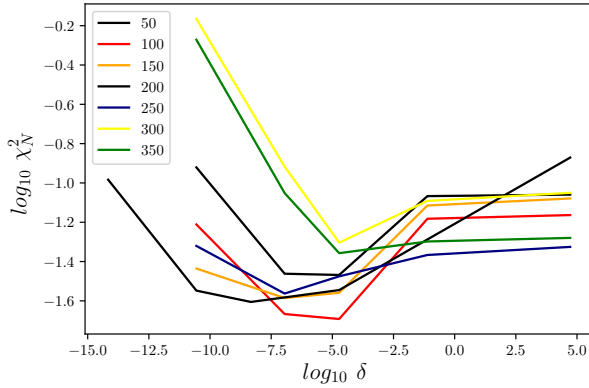


Fig. 3. $\log_{10}(\chi^2_N)$ versus $\log(\delta)$ for T_{eff} of a synthetic spectra having $T_{\text{eff}}=7600$ K, $\log g=2.50$ dex, $v \sin i=197 \text{ km s}^{-1}$, and with a S/N of 196. Each curve displays the minimization using different sets of nearest neighbor generated iteratively.

4.2 Inversion of simulated A stars

We used the LDB of Gebran et al. (2016) in which the effective temperature of the data varies from 6800 up to 11000 K. The wavelength region was chosen between 4450–4990 Å. This wavelength region harbors lines that are sensitive to all stellar parameters, and insensitive to microturbulent velocity which was adopted to be $\xi_t=2$ km/s based on the work of Gebran et al. (2014, 2016). The adopted resolution is 76000 as it corresponds to most of the analyzed stars in Sec. 5. The ranges of all parameters in the A-star LDB are summarized in Tab. 1.

Table 1. Ranges of the parameters used for the calculation of the A and FGK synthetic spectra LDB's

Parms	A stars	F/G/K
T_{eff} (K)	[6 800, 11 000]	[4 000, 8 000]
$\log g$ (dex)	[2.0, 5.0]	[3.0, 5.0]
$[M/H]$ (dex)	[-2.0, 2.0]	[-1.0, 1.0]
$v \sin i$ (km s^{-1})	[0, 300]	[0, 100]
$\lambda/\Delta\lambda$	76 000	50 000

Noise added synthetic spectra were calculated to be used as simulated observations. Around 1500 spectra were calculated for A stars. To analyze the effect of the sampling on the RSIR technique, we have inverted these spectra using 3 different LDB's. For the same range in all the parameter only the step was modified in each database. As an example, in the LDB 1, T_{eff} has a step of 100 K, whereas in LDB's 2 and 3, the steps are 200 K and 400 K, respectively. The same was done for all parameters and the details about the steps are found in Tab. 2. The sampling of the $v \sin i$ in the LDB's is not constant and depends on the value of $v \sin i$ (Gebran et al., 2016).

To compare the results of the inversion of PCA NN and RSIR for 1500 spectra, we estimate the root mean square error for both techniques, defined as:

$$\Lambda = \sqrt{\frac{\sum_{i=0}^N (y_i^{(\text{true})} - y_i^{(\text{inv})})^2}{N}}, \quad (13)$$

Table 2. Result of inversion for A stars using 3 different LDB's with different steps on 1500 noisy synthetic spectra

Test	Parms	step	Λ_{RSIR}	Λ_{PCA}
1	T_{eff} (K)	100	100.3	132
	$\log g$ (dex)	0.1	0.12	0.133
	$[M/H]$ (dex)	0.1	0.058	0.066
	$v \sin i$ (Km/s)	2: [0-20[5: [20-40[10: [40:300]	5.92	6.47
2	T_{eff} (K)	200	108	190
	$\log g$ (dex)	0.2	0.109	0.145
	$[M/H]$ (dex)	0.2	0.072	0.107
	$v \sin i$ (Km/s)	10	9.63	10.25
3	T_{eff} (K)	400	174	295
	$\log g$ (dex)	0.4	0.17	0.224
	$[M/H]$ (dex)	0.3	0.081	0.113
	$v \sin i$ (Km/s)	20	12.05	11.23

where $y_i^{(\text{true})}$ is the known parameter of the i^{th} synthetic spectrum and $y_i^{(\text{inv})}$ its corresponding inverted one.

The last 2 columns of Tab. 2 display the Λ results using the PCA NN search and the RSIR for the 3 LDB's. Comparing the Λ values of each approach, an improvement is achieved using RSIR of all parameters. One exception exists in the case of $v \sin i$ for test 3. The large $v \sin i$ step of the original LDB causes the PCA NN pre-processing stage to select inaccurate NN's. For most cases RSIR with a coarse sampling in parameters is producing more accurate inversions compared to PCA with a denser sampling. This directly infers a gain in computational time as a coarse sampling leads to smaller LDB.

To analyse the effect of the S/N on the inversions, we display in Fig. 4 the inverted T_{eff} as a function of the real T_{eff} for the 1500 A star spectra, for different S/N and different LDB's (tests 1, 2 and 3). The results of the PCA NN is affected both by the sampling size and the S/N of the analyzed stars, whereas for RSIR, with the pre-processing of PCA and a Tikhonov regularization this effects becomes less significant on the accuracy of the inversion.

4.3 Inversion of Simulated FGK type stars

The same procedure was applied to FGK star-like spectra. Around 2500 noisy spectra were produced in the ranges described Tab. 1. The chosen resolution of 50 000 is the same used in Paletou et al. (2015a). The wave-

Table 3. Result of inversion using 3 different LDB's with different steps on 2500 noisy synthetic FGK spectra

Test	parms	step	Λ_{RSIR}	Λ_{PCA}
1	T_{eff} (K)	100	106	117
	$\log g$ (dex)	0.1	0.121	0.136
	$[M/H]$ (dex)	0.1	0.061	0.063
	$v \sin i$ (Km/s)	2: [0-20[5: [20-40[10: [40-100]	1.72	2.88
2	T_{eff} (K)	200	109	236
	$\log g$ (dex)	0.2	0.132	0.289
	$[M/H]$ (dex)	0.2	0.066	0.115
	$v \sin i$ (Km/s)	2: [0-20[10: [20-40[10: [40-100]	1.88	3.17
3	T_{eff} (K)	400	127	368
	$\log g$ (dex)	0.4	0.147	0.45
	$[M/H]$ (dex)	0.4	0.07	0.17
	$v \sin i$ (Km/s)	4: [0-20[10: [20-40[10: [40-100]	2.25	5.243

length range was selected from 5000-5400 Å containing the MgI b triplet, a good indicator of $\log g$ and sensitive as well to T_{eff} . The microturbulent velocity was set to $\xi_t \sim 1 \text{ km s}^{-1}$ (Gebran et al., 2014). The inversion results as a function of the sampling steps are shown in Tab. 3. These results show similar behavior to that of the A stars in terms of improvement in accuracy while comparing RSIR to PCA NN search. In Fig. 4 we also over plot the T_{eff} for our F/G/K noisy synthetic spectra. The effect of inversion as function of S/N and sampling is very similar to the one of A stars.

5 Application to observed spectra

The performance of the RSIR has been tested on two samples of stellar spectra. The first sample is the one of the Spectroscopic Survey of Stars in the Solar Neighbourhood (S⁴N, Allende Prieto et al. 2004). These are spectra of bright FGK stars that are at distance less than 15 pc. We have estimated the S/N of these spectra in the wavelength range used for the inversion of the parameters [5000–5400Å]. This ratio ranges between 40 and 450. These spectra are at a resolution of $\lambda/\Delta\lambda \sim 50\,000$.

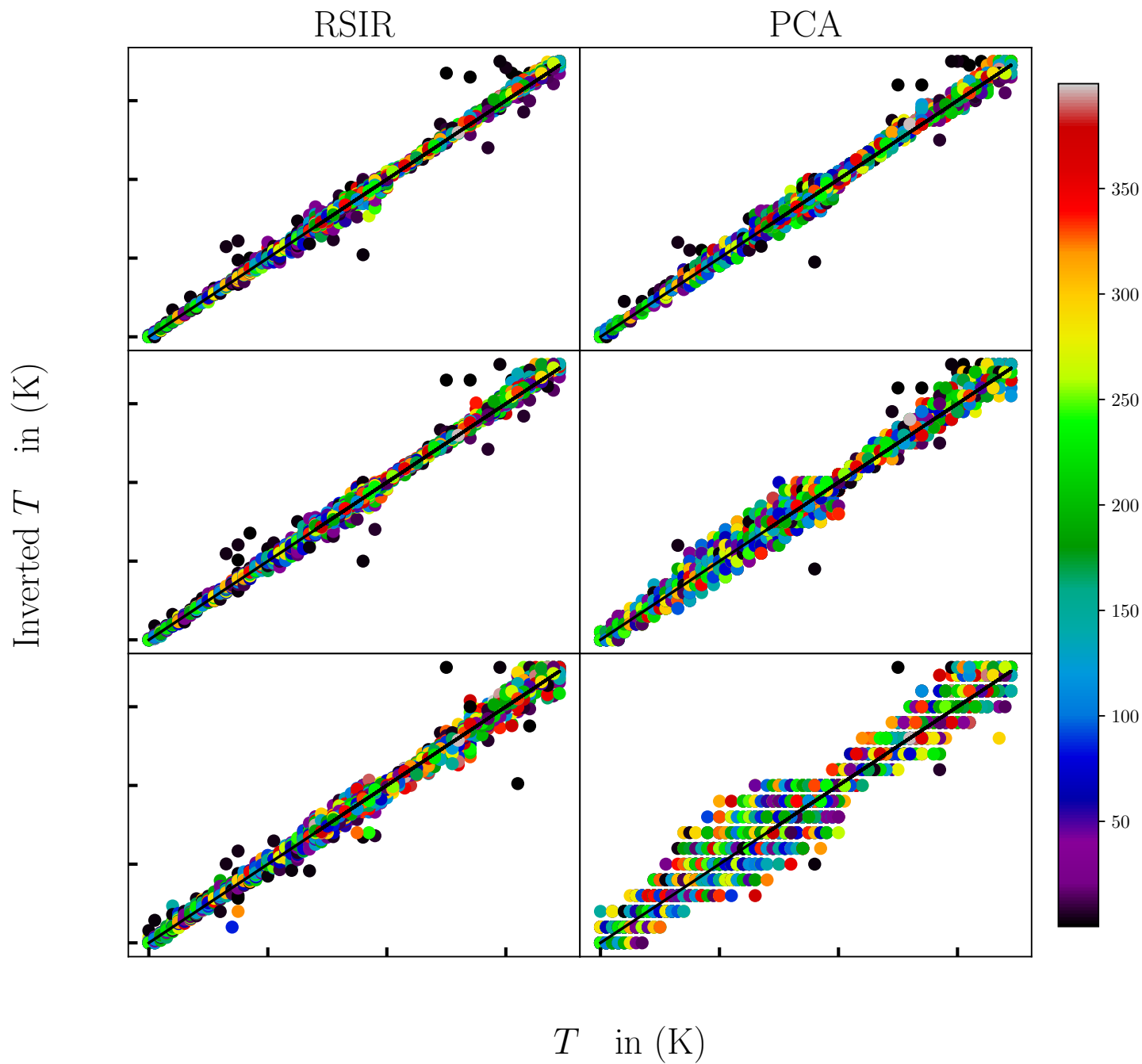


Fig. 4. Results of inversion of A to K type simulated stars. The comparison between PCA nearest neighbour search and RSIR based on two different steps implies an improvement.

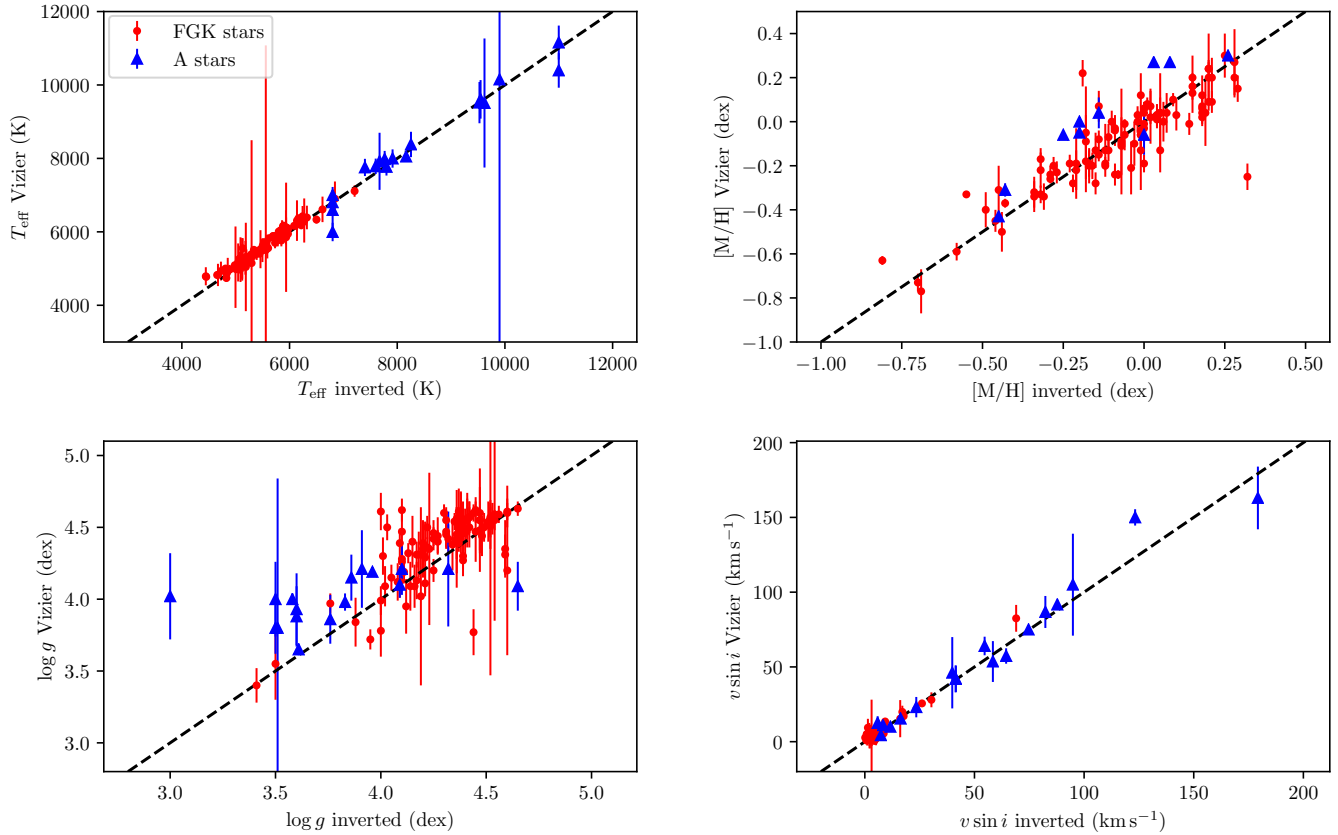


Fig. 5. Comparison between the inverted parameters for our FGK stars sample (filled circles) and A stars (filled triangle).

All the details about the acquisition and the reduction procedure of the S⁴N data can be found in Allende Prieto et al. (2004). These spectra were inverted using the database of Paleou et al. (2015a), made of 905 spectra retrieved from the ELODIE stellar library (Prugniel & Soubiran, 2001, Prugniel et al., 2007). We have used the MgI b triplet wavelength range as explained in Sec. 4.3. We then compared the values of the inverted parameters with the ones of Allende Prieto et al. (2004) and to the medians found in the Vizier catalog⁴ for all these stars. The main reason for using this catalog for our comparison is the necessity for reliable and objective catalogs which are constructed based on previous adopted values by the astronomical community.

Comparing our inverted T_{eff} to the ones of Allende Prieto et al. (2004), we found an average signed difference of 2.09 K with standard deviation of 102 K. For $\log g$, the average signed difference and the standard deviation are both 0.15 dex. For $[M/H]$ and $v \sin i$, we found -0.06 ± 0.08 dex and -0.21 ± 1.89 km s⁻¹, respectively. If we compare our inverted values to the median of Vizier, we find -85 ± 110 K, -0.07 ± 0.16 dex, 0.01 ± 0.10 dex and -0.50 ± 2.25 km s⁻¹ as a signed mean difference and a standard deviation between the catalogued values and the inverted ones for T_{eff} , $\log g$, $[M/H]$ and $v \sin i$, respectively. Figure 5 displays in filled circles, for the four parameters, the comparison between our inverted values for the FGK observed spectra and the ones derived from Vizier. We have also assigned the catalogues values an error bar corresponding to the standard deviation of the dispersion in the catalogues values for each star.

⁴ The query was performed using the method described in Paleou & Zolotukhin (2014)

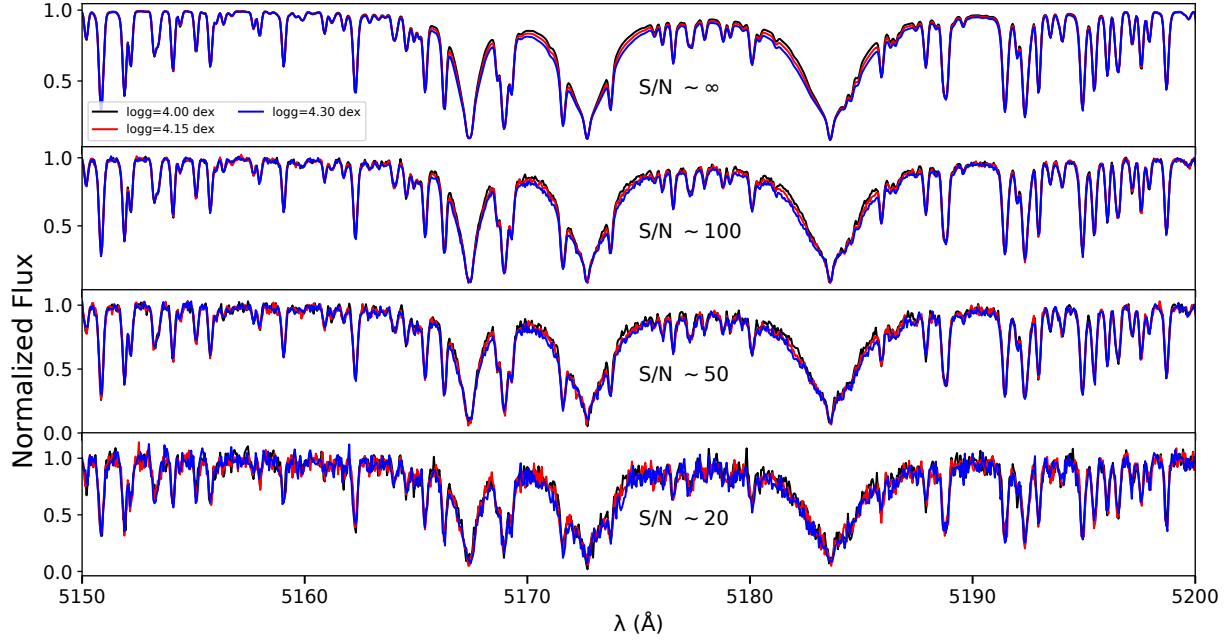


Fig. 6. Synthetic spectra of G stars having $[T_{\text{eff}}, \log g, [M/H], v \sin i]$ of 5200 K, [4.00, 4.15, 4.30 dex], 0.0 dex, 6 km s⁻¹. Each plot displays the spectra with different S/N.

The second sample of our analysis constitute of the well studied A stars of Gebran et al. (2016). These are the 19 stars that have been studied extensively by different authors using different techniques (Vega, Sirius A, HD 22484, HD 15318, HD 76644, HD 49933, HD 214994, HD 214923, HD 113139, HD 114330, HD 27819, HD 5448, HD 33256, HD 29388, HD 91480, HD 30210, HD 32301, HD 28355, and HD 222603) and have more than 120 references each. The source of these high resolution spectra is explained in detail in Gebran et al. (2016). They were observed using ELODIE, NARVAL, ESPaDOnS and SOPHIE spectrographs. ELODIE has a resolution of 42 000 whereas NARVAL, ESPaDOnS and SOPHIE are at a resolution of $\sim 76\,000$.

We have applied the RSIR on these data using the database of Test 1 in Sec. 4.2 at both resolutions. The S/N of these spectra is between 180 and 360. The inverted parameters of each star were compared to the ones retrieved from VizieR and added to the plots of Fig. 5 as filled triangles. We found an average signed difference and a standard deviation of -0.14 ± 245 K, -0.20 ± 0.30 dex, -0.11 ± 0.09 dex and -2.07 ± 8.5 km s⁻¹ for T_{eff} , $\log g$, $[M/H]$ and $v \sin i$, respectively between the inverted and the VizieR parameters.

Table 4. Estimation of the errors on the derived parameters for FGK and A stars.

Parameter	σ_{FGK}	σ_{A}
T_{eff} (K)	110	245
$\log g$ (dex)	0.16	0.30
$[M/H]$ (dex)	0.10	0.09
$v \sin i$ (km s ⁻¹)	2.25	8.50

5.1 The case of $\log g$

These results show that most of our inverted parameters are in agreement with previous studies. Considering that the most accurate parameters of these A and FGK stars are the VizieR median, our values are less spread with respect to the median than the ones of Paletou et al. (2015a) and Gebran et al. (2016). The standard deviations that we found could be assigned as an estimation of the errors on the derived parameters. We can therefore assign an accuracy of 110 K, 0.16 dex, 0.10 dex and 2.25 km s⁻¹, for T_{eff} , $\log g$, $[M/H]$, and $v \sin i$, respectively for FGK stars. For A stars, we found accuracies of 245 K, 0.30 dex, 0.09 dex, and 8.50 km s⁻¹, for T_{eff} , $\log g$, $[M/H]$, and $v \sin i$, respectively. These errors are summarized in Tab.4.

Surface gravity is systematically the most difficult parameter to determine, with typical errors of the order of 0.15 to 0.3 dex. Spectroscopic determinations of surface

gravity have been always assigned moderately large error bars, especially for A stars (Smalley, 2005). The same applies to FGK stars but with smaller error bars. Asteroseismic $\log g$ determinations remain the best tools for achieving accuracies less than 0.05 dex (Chaplin et al., 2014, Creevey et al., 2013, Hekker et al., 2013). RSIR is mainly based on finding the best set of spectra in the database that correspond to the observed one. As it is a spectroscopic method, we should not expect an accurate recovery for $\log g$. Using our values for $\Delta \log g$ we can, a posteriori figure out what it means in terms of discernibly between two spectra whose respective $\log g$ differ from this quantity. This also gives us relevant information about (i) which specific bandwidth(s) are the most sensitive to such differences, and (ii) how significant they are for various S/N.

Figures 6 and 7 display the variation in the spectrum profile as a function of $\log g$, fixing all the remaining parameters, for FGK and A stars, respectively. In Fig. 6, we calculated synthetic spectra for a typical G star with a T_{eff} of 5200 K, $[M/H]$ of 0.0 dex, $v \sin i$ of 6 km s^{-1} at a resolution of 50 000 and in the wavelength range of 5000–5400 Å. The only parameter that differs between the 3 spectra is $\log g$, ranging between 4.00 dex and 4.30 dex with a step of 0.15 dex. The upper panel of Fig. 6 displays the normalized flux of the synthetic spectra in the Mg I b triplet region. The flux level in this region is very sensitive to variation in $\log g$. The following panels displays the same spectra for different values of S/N. When no noise is added (panel with $S/N \sim \infty$), the distinction between the 3 spectra is clear but when S/N starts to decrease, the distinction between the noisy spectra becomes harder to detect. This shows that for a S/N in the order of 100, the noisy spectra with $\log g$ of 4.00 and 4.15 dex are very similar and therefore the best corresponding synthetic spectrum in our LDB could have a $\log g$ varying at least 0.15 dex from the correct value. We are not trying to quantify the minimum S/N required for an accurate inversion of $\log g$ as the RSIR is not based on a pixel-to-pixel comparison, but we are showing the effect of our derived standard deviations in $\log g$ on the flux for noisy spectra. Figure 7 displays a similar behaviour for A stars having similar T_{eff} of 8500 K, $[M/H]$ of 0.0 dex, $v \sin i$ of 40 km s^{-1} , at a resolution of 76 000 in the wavelength range of 4500–5000 Å. Surface gravity of these spectra ranges between 3.60 and 4.20 dex with a step of 0.30 dex. This figure shows a similar behavior to that of Fig. 6. At a S/N of ~ 150 , the distinction between spectra having a difference of 0.30 dex in $\log g$, becomes hardly noticeable. Figures 6 and 7 also show that the effect of weak metallic lines, on the derivation of $\log g$, becomes negligible as

the S/N decreases. The $\log g$ information that these lines contain is mainly lost in the noise.

6 Discussion and conclusion

RSIR tests for nearly 4000 synthetic stars of different spectral type and different noise levels showed an improvement over the PCA-based method of Paletou et al. (2015a,b) and Gebran et al. (2016) for the inversion of stellar parameters. Results of Tabs. 2 and 3 and Fig. 4, for FGK and A stars, show that for most of the tests, the Λ values of RSIR are lower than nearest neighbor PCA approach. Having a prior information about the star using PCA as a pre-process allows us to narrow down the selection of the an optimized reduced databases. This decreases drastically the size of the LDB's. Achieving lower Λ with bigger steps helps in decreasing the prohibitive computation time for the construction of databases and the calculations of the PC's. Simulated tests revealed that computation time of RSIR is nearly 1% of that of the process of PCA nearest neighbor approach.

One should be very careful while increasing the size of steps of the parameters because the PCA pre-processing step could deviates drastically from the true inverted values therefore excluding the spectra that actually best describes the observed ones.

Application to observed FGK and A stars reveal a good agreement between the inverted parameters and the ones derived in previous studies. The comparison with Vizier catalog values show an improvement in the derived parameters as compared to the results of Paletou et al. (2015a) and Gebran et al. (2016) for the same stars and LDB's. Surface gravity remains the parameter with the least accuracy. Our derived errors on $\log g$ are in the order of 0.15–0.30 dex. Smarter LDB's should be therefore considered, say, “adaptive sampling” (in the parameters under study), taking care with more caution of the flux typical variations at the most sensitive wavelength (sub-)domains, *together with* the S/N of the observations, instead of the a priori sampling in the parameters. Also, a commonly reported issue with the inversion of stellar parameters using a LDB of synthetic spectra are the so-called “ambiguities”. This means that two sets of distinct parameters may generate spectra which are beyond “discernibility”. Given a set of observed spectra to characterize, we could naturally relate that discernibility to their level of S/N. Using a nearest neighbour search PCA-based method, for instance, such a level of S/N can easily be translated into a threshold of distance δ_{PCA} . Then,

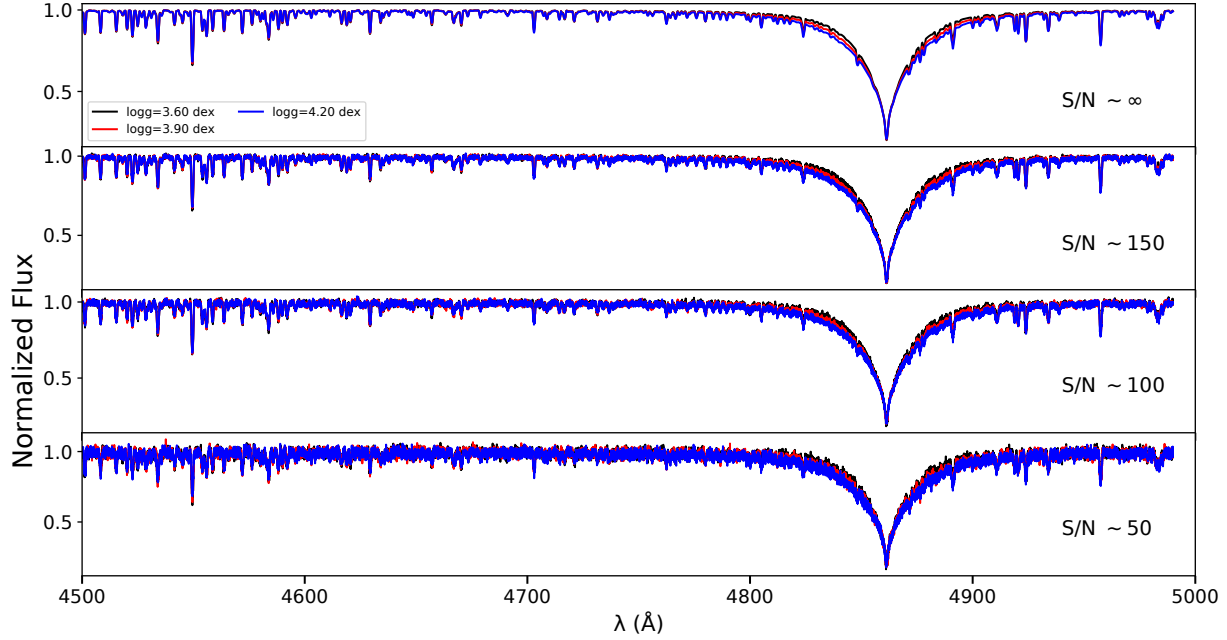


Fig. 7. Same as for Fig. 6 but for A stars having having $[T_{\text{eff}}, \log g, [M/H], v \sin i]$ of $[8\,500\text{ K}, [3.60, 3.90, 4.20\text{ dex}], 0.0\text{ dex}, 40\text{ km s}^{-1}]$.

we can anticipate that, instead of relying on LDB's usually made using a priori sampling *in the parameters*, a smarter DB should rely on δ_{PCA} instead. This would imply to set up LDB's for stellar fundamental parameters in a radically different fashion vs. common practices. In the frame of PCA, it would be more relevant to sample properly the full range of parameters with a “constrained-random” process ensuring that there are no nearest neighbours closer than δ_{PCA} . Such a “sieve algorithm” was first proposed by López Ariste & Casini (2002) in the context of the characterization of magnetic fields from spectropolarimetric data (see also Casini et al. 2013).

Available online databases are usually calculated with large steps in T_{eff} and $\log g$. Our RSIR technique, as it does not require small steps in the LDB, is a good tool to be used with online available synthetic spectra such as the POLLUX⁵ database (Palacios et al., 2010) that contains stars with temperature ranging between 3 000 and 50 000 K or TLUSTY Non-LTE Line-blanketed Model Atmospheres of O-Type Stars (Lanz & Hubeny, 2003) with T_{eff} ranging between 27 500 and 55 000 K with 2 500 K steps, and $\log g$ between 3.0 and 4.75 with steps of 0.25 dex. We can also mention the PHOENIX (Husser et al., 2013) models database for stars having $T_{\text{eff}} < 12\,000\text{ K}$ and

the AMBRE (de Laverny et al., 2012) project that contains high-resolution FGKM stellar synthetic spectra.

As an output of the new Gaia Data Release 2, Cropper et al. (2018) describe the Gaia RVS specification as well as the predicted performance at the end of the mission. Gaia RVS will provide us with a large number of spectra in the calcium triplet regime (845–872 nm). This triplet is very sensitive to T_{eff} and $\log g$. The medium resolution (11 500) of the RVS and the small range in wavelength would require LDB smaller than the ones used in our work, leading to a fast application of the RSIR. As we did for the inversion of the S⁴N data in Sec. 5, LDB could be constructed with real observed stars having well known fundamental parameters and with the same resolution. Finally, since RSIR is based on single parameter inversion process, one can also incorporate other parameters if there is a computational power facilitating this step and applicable theoretical models, for example, individual chemical abundances.

References

- Abolfathi, B., Aguado, D. S., Aguilar, G., et al. 2017, arXiv:1707.09322 .
Data: http://www.sdss.org/dr14/data_access/volume/

- Allende Prieto, C., Barklem, P. S., Lambert, D. L., & Cunha, K. 2004, *A & A*, 420, 183
- Alves, S., Benamati, L., Santos, N. C., et al. 2015, *MNRAS*, 448, 2749
- Bailer-Jones, C. A. L., Irwin, M., & von Hippel, T. 1998, *MNRAS*, 298, 361
- Bernard-Michel, C. and Douté, S. and Fauvel, M. and Gardes, L. and Girard, S. 2007.[Research Report], INRIA(inria-00187444v2), 91
- Bernard-Michel, C., Douté, S., Fauvel, M., Gardes, L., & Girard, S. 2009, *Journal of Geophysical Research (Planets)*, 114, E06005
- Boeche, C., Smith, M. C., Grebel, E. K., et al. 2018, *AJ*, 155, 181
- Buchhave, L. A., Latham, D. W., Johansen, A., et al. 2012, *Nature*, 486, 375
- Casini, R., Asensio Ramos, A., Lites, B. W., & López Ariste, A. 2013, *APJ*, 773, 180
- Chaplin, W. J., Basu, S., Huber, D., et al. 2014, *ApJS*, 210, 1
- Casey, A. R., Hogg, D. W., Ness, M., et al. 2016, *arXiv:1603.03040*
- Castelli, F., & Kurucz, R. L. 2003, *Modelling of Stellar Atmospheres*, 210, A20
- Cayrel, G., & Cayrel, R. 1963, *APJ*, 137, 431
- Cayrel de Strobel, G. 1969, in *Proc. of the 3rd Harvard-Smithsonian Conf. on Stellar Atmospheres*, ed. O. Gingerich, 35
- Creevey, O. L., Thévenin, F., Basu, S., et al. 2013, *MNRAS*, 431, 2419
- Cropper, M., Katz, D., Sartoretti, P., et al. 2018, *arXiv:astro-ph/1804.09369*
- Cui, X.-Q., Zhao, Y.-H., Chu, Y.-Q., et al. 2012, *Research in Astronomy and Astrophysics*, 12, 1197.
Data: <http://dr5.lamost.org/>
- Dieterich, S., Henry, T. J., Benedict, G. F., et al. 2017, *American Astronomical Society Meeting Abstracts #229*, 229, 240.30
- Fabbro, S., Venn, K. A., O'Brian, T., et al. 2018, *MNRAS*, 475, 2978
- Gebran, M., Monier, R., Royer, F., Lobel, A., & Blomme, R. 2014, *Putting A Stars into Context: Evolution, Environment, and Related Stars*, 193
- Gebran, M., Farah, W., Paletou, F., Monier, R., & Watson, V. 2016, *A & A*, 589, A83
- Gill, S., Maxted, P. F. L., & Smalley, B. 2018, *arXiv:1801.06106*
- Hekker, S., Elsworth, Y., Mosser, B., et al. 2013, *A & A*, 556, A59
- Hubeny, I., & Lanz, T. 1992, *A & A*, 262, 501
- Husser, T.-O., Wende-von Berg, S., Dreizler, S., et al. 2013, *A & A*, 553, A6
- Jolliffe, I. T. 1986, *Springer Series in Statistics*, Berlin: Springer, 1986
- Katz, D., & Brown, A. G. A. 2017, *SF2A-2017: Proceedings of the Annual meeting of the French Society of Astronomy and Astrophysics*, 259
- Kiefer, J. 1953, *Proceedings of the American Mathematical Society*, 4, 502
- Kreyszig, E. , *Advanced Engineering Mathematics* 2010. (John Wiley & Sons) p .864-871
- Kurucz, R. L. 1992, *RMXAA*, 23.
- Lanz, T., & Hubeny, I. 2003, *ApJS*, 146, 417
- Latham, D. W., Stefanik, R. P., Torres, G., et al. 2002, *AJ*, 124, 1144
- de Laverny, P., Recio-Blanco, A., Worley, C. C., & Plez, B. 2012, *A & A*, 544, A126
- Li, K.C 1991, *Journal of the American Statistical Association*, 86, 414, 316
- López Ariste, A., & Casini, R. 2002, *APJ*, 575, 529
- McWilliam, A. 1990, *ApJS*, 74, 1075
- Morris, M., Kaiser, M. E., Bohlin, R., Kurucz, R., & ACCESS Team 2018, *American Astronomical Society Meeting Abstracts #231*, 231, #355.27
- Ness, M., Hogg, D. W., Rix, H.-W., Ho, A. Y. Q., & Zasowski, G. 2015, *APJ*, 808, 16
- O'Mullane, W., & LSST Data Management Team 2018, *American Astronomical Society Meeting Abstracts #231*, 231, #362.10
- Palacios, A., Gebran, M., Josselin, E., et al. 2010, *A & A*, 516, A13
- Paletou, F., & Zolotukhin, I. 2014, *arXiv:1408.7026*
- Paletou, F., Böhm, T., Watson, V., & Trouilhet, J.-F. 2015, *A & A*, 573, A67
- Paletou, F., Gebran, M., Houdebine, E. R., & Watson, V. 2015, *A & A*, 580, A78
- Perryman, M. A. C., de Boer, K. S., Gilmore, G., et al. 2001, *A & A*, 369, 339.
Data: <https://www.cosmos.esa.int/web/gaia/dr2>
- Prugniel, P., & Soubiran, C. 2001, *A & A*, 369, 1048
- Prugniel, P., Soubiran, C., Koleva, M., & Le Borgne, D. 2007, *arXiv:astro-ph/0703658*
- Re Fiorentin, P., Bailer-Jones, C. A. L., Lee, Y. S., et al. 2007, *A & A*, 467, 1373
- Schönrich, R., & Bergemann, M. 2014, *MNRAS*, 443, 698
- Shevlyakova, M., & Morgenthaler, S. 2014, *Statistical Papers*, 55, 1
- Smalley, B., Zverko, J., Žižňovský, J., Adelman, S.J., and Weiss, W.W 2004. *The A-Star Puzzle* . *Proc. IAU Symp. 224* (Cambridge Univ. Press Cambridge).
- Smalley, B. 2005, *Memorie della Societa Astronomica Italiana Supplementi*, 8, 130
- Stoehr, F., White, R., Smith, M., et al. 2008, *Astronomical Data Analysis Software and Systems XVII*, 394, 505
- Torres, G., Neuhäuser, R., & Guenther, E. W. 2002, *AJ*, 123, 1701
- Valenti, J. A., & Piskunov, N. 1996, *A & AS*, 118, 595
- C.R. Vogel. *Computational methods for inverse problems*. Society for Industrial and Applied Mathematics, 2002
- Watson, V., Trouilhet, J., Paletou, F., & Gebran, M. 2017, *arXiv:1706.10121*
- Wilkinson, D. M., Maraston, C., Goddard, D., Thomas, D., & Parikh, T. 2017, *MNRAS*, 472, 4297
- Xiang, M.-S., Liu, X.-W., Shi, J.-R., et al. 2017, *MNRAS*, 464, 3657

