

Java sur Raspberry Pi

Entretien avec Frank Delporte



C. J. Abate (Elektor)

Vous pouvez faire tourner Java sur Raspberry Pi. Dans cet entretien, l'auteur Frank Delporte parle des avantages de Java, ses projets basés sur Raspberry Pi et raconte comment il en est venu à marier la carpe et le lapin.

Faire tourner Java sur Raspberry Pi, ça vous dirait ? C'est ce que vous propose Frank Delporte. Son nouveau livre, *Getting Started with Java on the Raspberry Pi* [1], est une excellente ressource pour les programmeurs professionnels et les mordus avides d'expérimentation, désireux d'apprendre à leur propre rythme. Dans cet entretien, Delporte parle des avantages de la combinaison de Java et de Raspberry Pi, ainsi que de ses expériences en tant que programmeur.

Programmation, conception et rédaction

Abate : Félicitations pour la publication du livre *Getting Started with Java on the Raspberry Pi* (Elektor 2020). Nous reparlons du livre dans un instant, mais quel est ton sentiment maintenant que tous les travaux d'écriture et d'édition sont derrière toi ? As-tu aimé ce genre de travail ?

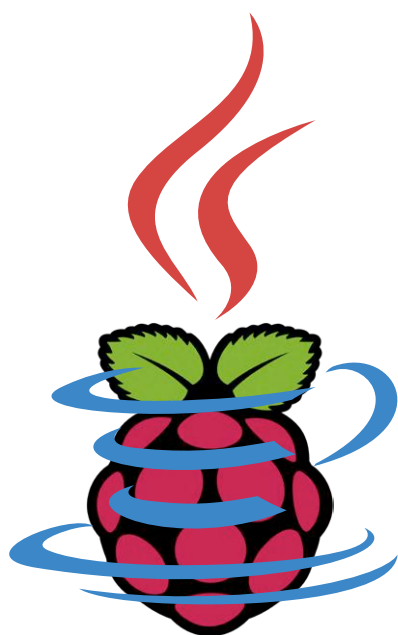
Delporte : J'aime écrire, il n'y a qu'à voir mon blog [2], mais écrire un livre entier, c'est du travail. Il faut l'écrire, bien sûr, mais avant il faut rassembler toutes les infos, faire des recherches, expérimenter, dessiner les schémas, mener des entretiens, et après il faut relire, etc. La satisfaction de tenir enfin entre mes mains mon premier livre imprimé (fig. 1) me récompense largement pour tout ce travail.

Abate : Nous avons présenté ton espace de travail à domicile [3] sur le site d'Elektor en mai 2020. Travailles-tu toujours à la maison en raison de la COVID-19 ?

Delporte : Ici la situation est redevenue presque normale. Je travaille toujours à la maison, mais pas pour cause de Covid. Pour me concentrer sur une tâche précise, rien de tel que de travailler chez moi.

Abate : Ton parcours est intéressant, à la fois concepteur de logiciels, responsable technique, auteur et monteur vidéo. Comment t'y es-tu pris ?

Delporte : J'ai toujours été intéressé par la technique et le fonctionnement des choses. J'étais le genre de gamin qui démonte chaque machine à café, radio ou tout autre appareil cassé. Je n'arrivais pas à les réparer toutes, mais chaque fois j'apprenais quelque chose de nouveau ! Adolescent, j'ai participé comme DJ à une émission de radio locale, ce qui m'a donné la possibilité d'expérimenter davantage l'électronique. C'est ainsi que j'ai décidé d'étudier dans une école (technique) de cinéma, pour apprendre les techniques du cinéma, de la radio et de la télévision, l'étalonnage des caméras et les connexions de tous ces appareils et les techniques d'enregistrement. Après mon diplôme, le montage informatique a bouleversé la production



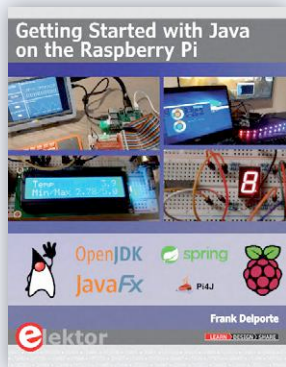


Figure 1. Le livre de Frank Delporte : *Getting Started with Java on the Raspberry Pi*.

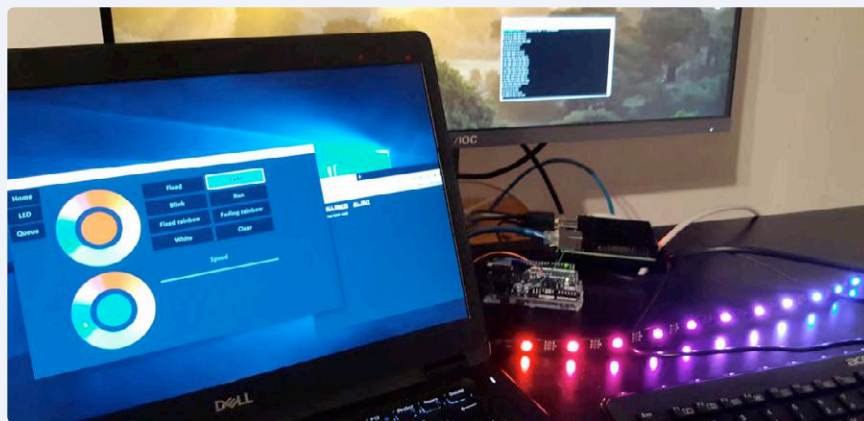


Figure 2. La configuration de la bande de LED.

vidéo, et c'est ainsi que je suis revenu à la programmation parce que des clients voulaient leur vidéo d'entreprise sur CD-ROM puis sur l'internet.

Abate : Dans ta biographie figure le Commodore 64. Te souviens-tu de tes premières expériences ?

Delporte : Je n'avais qu'un seul jeu sur mon C64, parce que ce qui m'intéresse, c'est la programmation. Grâce à Elektor, j'ai trouvé un livre (vers 1987) avec une carte électronique à huit relais que l'on pouvait commander en Basic sur le C64. Je l'ai utilisée pour commander mon train Lego ; des interrupteurs magnétiques reliés aux ports du joystick détectaient la position du train. Ce fut ma première combinaison réussie de logiciel et de matériel. Aujourd'hui, un tel projet serait beaucoup plus facile (et moins cher) avec un Arduino ou un Raspberry Pi et les nombreuses cartes d'extension géniales.

Abate : Quels étaient tes objectifs de carrière en 1994 à ta sortie du NARAFI (Nationaal Radio en Filmtechnisch Instituut) ?

Delporte : Mon premier travail a été le montage vidéo pour une TV locale, et après quelques années, puis j'ai fait ça comme pigiste. Sans objectif de carrière précis, je suis passé d'un emploi à l'autre en apprenant beaucoup en cours de route. C'est ainsi que je suis passé de monteur vidéo à développeur multimédia en créant des présentations d'entreprise sur CD-ROM puis DVD. Avec l'essor de l'internet, je suis passé au développement web car mes clients voulaient diffuser les mêmes informations sur des sites avec un système de gestion de contenu. Grâce à ces connaissances, je suis devenu développeur Java et responsable technique du développement de produits utilisant cette technologie.

Enseigner le codage

Abate : Quand as-tu commencé à organiser des sessions de CoderDojo ? Et quel genre de cours donnes-tu ?

Delporte : Dans chaque entreprise où j'ai travaillé, le défi a été de trouver les bons collègues sur le plan technique. Les techniques et certainement l'informatique sont encore trop masculines. Pour moi, l'ingénierie est magique. Quelques lignes de code ou quelques composants électroniques suffisent pour construire des "trucs". Pourquoi tant d'enfants, qui pourtant aiment construire et expérimenter, s'arrêtent-ils soudain de le faire et ne choisissent pas une orientation où ils peuvent continuer à inventer ? CoderDojo [4] est un club gratuit où des bénévoles aident les jeunes de 7 à 18 ans à expérimenter des *trucs numériques*. Nous programmons en Scratch (par blocs), construisons des mondes dans Minecraft avec JavaScript, contrôlons l'électronique avec Arduino, construisons des robots avec Lego, et bien davantage. En 2013, j'ai lancé un tel club à Ypres et à Roulers, en Belgique, et je dirige toujours celui d'Ypres. La crise sanitaire nous a contraints à suspendre provisoirement nos rencontres informelles à plusieurs autour d'un PC. Grâce au CoderDojo et à d'autres initiatives STEM (Science, Technologie, Ingénierie et Mathématiques), nous constatons une progression du nombre d'étudiants en technique, garçons et filles !

Abate : Quand t'es-tu dit : "Hé, je suis un bon formateur, et je peux aider d'autres personnes intéressées par Java" ? Ou est-ce un ami ou un collègue qui t'a indiqué la direction à suivre ?

Delporte : J'aime expliquer et je crois fermement au principe d'apprendre en apprenant. C'est ce que je fais au CoderDojo avec les jeunes, mais aussi sur mon blog

et au travail. Pour bien comprendre un sujet, il faut être capable de l'expliquer et vice versa. Les articles que j'écris pour mon blog naissent toujours de quelque chose que je veux essayer sans (encore) vraiment savoir comment le faire. Au fil du processus de compréhension, j'écris les étapes suivies et ce qui a fonctionné et ce qui a cafouillé. C'est ainsi que j'apprends et que j'acquiers des connaissances à partager avec d'autres.

Zoom sur Java

Abate : Raconte-nous ton histoire avec Java. L'as-tu appris par curiosité ? Était-ce pour un cours ? Ou pour le travail ?

Delporte : Pour produire des applications multimédias, j'ai dû apprendre *ActionScript* (et même *Lingo* avant). Plus tard, je suis passé au C# et au SQL pour les applications web. On aura compris que j'apprends surtout en expérimentant, mais aussi en lisant des livres et en suivant de courts cours (en ligne). Lorsque j'ai commencé à *Televis Rail* en 2010, j'ai rejoint une équipe qui utilisait déjà Java. Le passage de C# à Java a été très facile. Après toutes ces années de programmation, je réalise que c'est avec mes collègues que j'apprends le plus ! Partager son travail avec d'autres lors de présentations, améliorer le code avec des demandes d'extraction, accepter des commentaires comme moyen d'amélioration sont autant d'excellents moyens d'apprendre des autres.

Depuis que j'ai commencé à expérimenter en Java sur Raspberry Pi, je me suis impliqué dans des projets et des discussions sur les logiciels libres, et c'est un monde complètement nouveau pour moi, où je rencontre beaucoup de gens brillants, également prêts à partager leurs connaissances et leur expérience. Il ne passe pas un jour sans que je sois étonné par tout ce



Figure 3. Les composants installés de la console de commande de la cabine de batterie.

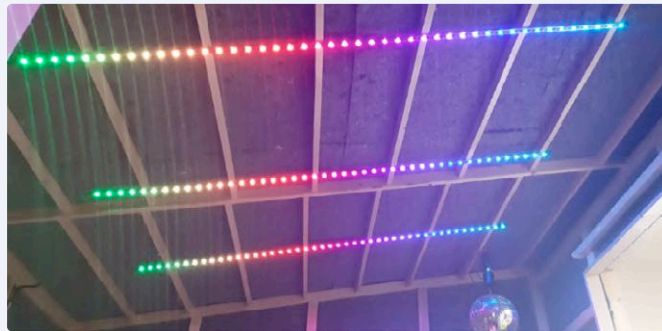


Figure 4. Effet d'arc-en-ciel réalisé avec des bandes de LED.



Java n'est pas seulement un langage de programmation, c'est aussi une machine virtuelle qui exécute le code Java.



que ces projets et ces personnes peuvent nous apprendre. Vous n'êtes pas obligé de contribuer au code, mais vous pouvez aussi rejoindre un tel projet en examinant les demandes de retrait, en aidant à tester ou à documenter le code.

Abate : Es-tu anti-python ou anti-C ? Je suppose que non, mais je ne peux pas ne pas te poser la question.

Delporte : Certainement pas, je me garde de ce qui est « anti ». Il n'y a pas de mauvais langages de programmation ! On ne devrait jamais perdre de vue que le meilleur outil pour faire un travail est celui que l'on connaît le mieux. Dans mon cas, c'est Java et JavaFX si je veux faire une application avec une belle interface utilisateur. Dans mon livre, j'ai aussi utilisé Python pour commander un affichage de nombres par LED et un Arduino avec des bandes de LED. Pour chaque projet (fig. 2), il faut décider quel est le meilleur outil, le meilleur langage de programmation ou la meilleure plateforme. Et une fois ta décision prise, vas-y ! Plus tard, tu verras que ce n'était peut-être pas le meilleur choix, mais tu auras tant appris que tu ne regretteras rien.

Travailler avec Java sur le Pi

Abate : Tu bloques depuis 2007 sur les sujets techniques. Ton premier post sur Raspberry Pi semble être "Pong on a Raspberry Pi" [5] en décembre 2017. Parle-nous de tes premières expériences avec RPi ? Quand as-tu commencé ?

Delporte : Quand j'ai commencé avec CoderDojo, certains moniteurs avaient

déjà une expérience avec Arduino et Raspberry Pi. Ils sont venus avec leurs kits au club. La puissance de ces cartes bon marché m'a stupéfié quand j'ai compris tout ce qu'on peut en tirer en les combinant avec quelques composants électroniques.

Je bloguais déjà depuis un certain temps, mais mon premier projet Raspberry Pi "public" était en effet ce jeu de Pong que nous avons utilisé pour certaines activités de l'école de mon fils. J'ai utilisé Python pour l'interface utilisateur ; mais pour tout dire, sans grand plaisir. Pour ce genre d'application, je préfère JavaFX pour lequel il existe même un très beau cadre de jeu : FXGL [6].

Abate : As-tu chez toi ou dans ton espace de travail des applications basées sur le RPi qui fonctionnent ?

Delporte : J'ai commencé avec Java sur le Raspberry Pi pour construire un contrôleur de batterie [7] pour mon fils. Il s'agit d'une interface utilisateur à écran tactile permettant de commander plusieurs lumières à l'aide d'une carte à relais et des bandes de LED commandées par Arduino (fig. 3).

J'ai ainsi appris à utiliser la communication série entre les deux cartes et I²C pour commander les relais. Dans mon livre, j'ai étendu cette méthode et j'ai utilisé une file d'attente Mosquitto pour échanger des messages entre d'autres cartes et PC.

Abate : Sur quoi d'autre travailles-tu ces jours-ci ? De nouveaux projets, programmes ou livres ?

Delporte : Je poursuis mes expériences avec Java sur le Raspberry Pi, bien sûr. J'ai écrit d'autres articles de blog sur ce sujet et j'ai également expérimenté d'autres technologies Java (Quarkus [8], Spring, GraalVM) et des systèmes d'exploitation à 64 bits sur le forum.

J'ai également rejoint l'équipe Pi4J. Pi4J est un cadre et une bibliothèque permettant de combiner des applications Java avec toute la puissance des GPIO du Raspberry Pi. Ce projet a été lancé par Robert Savage qui cherchait des membres pour l'équipe afin de le porter à une nouvelle génération qui supporte entièrement Java 11+ et le Raspberry Pi 4 avec des modules Java et une architecture facilement extensible. Je suis enthousiaste à propos de la deuxième version de ce cadre, que nous espérons publier bientôt.

Abate : Revenons au livre, *Getting Started with Java on the Raspberry Pi*. Pourquoi l'as-tu écrit ?

Delporte : Pour le projet de la cabine de batterie, j'ai dû découvrir comment utiliser Java sur le Pi, comment installer la bonne version de JavaFX, comment contrôler le GPIO et un Arduino, etc. C'est alors que j'ai écrit mon premier article [9] publié dans MagPi (juillet 2019 [10]).

Elektor m'a demandé si cela pouvait faire l'objet d'un livre. Comme je ne trouvais pas de livre récent sur ce sujet et que Java a connu de grands changements ces dernières années, cette question m'a interpellé et dès le lendemain, j'ai commencé à écrire. Cela m'a pris plus de six mois et

beaucoup de soirées et de nuits, mais je me suis bien amusé en écrivant et en expérimentant. Pourvu que ce soit aussi amusant de lire le livre et d'essayer mes projets !

Abate : As-tu des conseils à donner aux ingénieurs ou aux mordus qui envisagent d'utiliser Java pour leurs projets sur Raspberry Pi ?

Delporte : Essayez ! Java est toujours l'un des meilleurs langages de programmation au monde. Que vous soyez un développeur Java expérimenté ou que vous partiez de zéro, il y a beaucoup à apprendre et à expérimenter lorsque vous combinez Java avec un Raspberry Pi et des composants électroniques.

Les exemples de mon livre utilisent des composants très bon marché. Ils ne sont pas difficiles à trouver. Tous les exemples du livre peuvent être utilisés pour donner vie aux projets de vos propres rêves. Le circuit de commande conçu pour la cabine de batterie de mon fils est une combinaison de plusieurs de ces exemples (fig. 4).

Abate : Quels ont été les retours d'information jusqu'à présent ?

Delporte : Python est évidemment le premier langage sur le Raspberry Pi (mais oui, c'est de là que vient le Pi) au point que certains le considèrent comme le seul bon choix, mais j'ai reçu beaucoup de commentaires positifs et nombre de questions sur ce sujet. J'ai même eu la chance d'écrire un billet pour l'*Oracle Java Magazine* [12], très lu et beaucoup partagé ! Il existe un

intérêt évident pour ce sujet, et la future nouvelle génération de Pi4J facilitera encore la création d'applications puissantes.

Abate : Y a-t-il un langage de programmation que tu ne connais pas, mais que tu comptes apprendre ? Y a-t-il un matériel que tu envisages d'essayer ?

Delporte : Java n'est pas seulement un langage de programmation, mais aussi une machine virtuelle qui exécute le code Java. Sur cette même VM, vous pouvez également exécuter Scala, Kotlin et bien d'autres langages. Il y a donc encore beaucoup de choses à explorer dans ce monde. Pour le projet Pi4J, je veux étendre le site web du code d'exemple et de la documentation, je devrai donc mettre en place plusieurs petits exemples de matériel et apprendre beaucoup de nouvelles choses moi-même.

Succès de la programmation

Abate : Concluons sur ton plus grand succès en matière de conception ou de programmation. Y a-t-il un projet spécifique (logiciel ou matériel) qui se démarque ? Qu'y avait-il de difficile dans ce projet ? Qu'as-tu appris ?

Delporte : Dans mon travail à Televic, nous utilisons une combinaison de Java et de programmation intégrée pour connecter plusieurs serveurs et sources de données afin d'afficher en temps réel les informations pour les passagers sur les écrans de trains entiers. Un défi technique ! Parcourir un train en marche avec 100 écrans qui affichent en temps réel les trains en partance dans la prochaine gare avec les retards et les correspondances, c'est le pied !

Le flux nécessaire pour acheminer toutes ces données par des connexions sans fil peu fiables (les signaux GSM ne sont pas vraiment conçus pour les trains à grande vitesse) est un véritable chef-d'œuvre dont je suis très fier et que nous avons pu réaliser avec une petite équipe. Et je suis tout aussi impressionné par les enfants du CoderDojo qui ont réussi à réaliser leur premier jeu *Flappy Bird* en Scratch ou à faire clignoter une LED avec Arduino ! 

200503-03

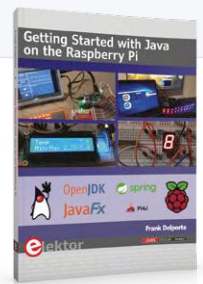


PRODUITS

➤ Livre :

Getting Started with Java on the Raspberry Pi par **Frank Delporte**

www.elektor.fr/19292



LIENS

- [1] **Livre :** *Getting Started with Java on the Raspberry Pi* : www.elektor.fr/getting-started-with-java-on-the-raspberry-pi
- [2] **WebTechie :** blog de l'auteur: <http://webtechie.be/>
- [3] **"A Software Developer's Space for DIY Projects and Writing" :** www.elektormagazine.com/news/electronics-workspace-software-developers-space
- [4] **CoderDojo :** <http://coderdojo.com/>
- [5] **"Pong on a Raspberry Pi" :** <http://webtechie.be/post/2017-12-20-pong-on-a-raspberry-pi/>
- [6] **"Getting Started with FXGL Game Development" :** <http://webtechie.be/post/2020-05-07-getting-started-with-fxgl/>
- [7] **"Drumbooth Controller with Raspberry Pi and JavaFX" :** <http://webtechie.be/post/2020-03-30-drumbooth-controller-with-java-javafx-raspberrypi-arduino/>
- [8] **Quarkus :** <http://webtechie.be/post/2020-07-28-spring-versus-quarkus-rest-h2-db-on-raspberry-pi/>
- [9] **WebTechie Articles :** <http://webtechie.be/articles/>
- [11] **MagPi (en français !)** : www.magpi.fr/
- [12] **"Getting Started with JavaFX on Raspberry Pi" :** <http://blogs.oracle.com/javamagazine/getting-started-with-javafx-on-raspberry-pi>



analyse de données et intelligence artificielle en Python

Interpréter les données réelles avec NumPy, pandas et le scikit-learn

Angelo Cardellicchio (Italie)

L'analyse et l'interprétation des données provenant de notre environnement est un sujet d'intérêt croissant. Ces données font désormais partie intégrante de notre vie quotidienne, qu'il s'agisse des données climatiques ou des données acquises au cours de processus de fabrication intelligents. La quantité de données nous permet, en théorie, de caractériser n'importe quel phénomène. Cependant il faut pour cela de nombreuses compétences, théoriques et pratiques. Nous examinons ici certaines de ces techniques et utilisons Python pour l'analyse de données du monde réel.

Des termes tels que *big data* (= données massives) et *intelligence artificielle* appartiennent au langage quotidien. Cela est principalement dû à deux facteurs. Le premier est la diffusion croissante et omniprésente des systèmes d'acquisition de données qui a permis la création de *référentiels de connaissances* pratiquement illimités. Le second est la croissance continue des capacités de calcul, grâce à l'utilisation généralisée des GPGPU (unités de traitement graphique polyvalentes) [1], qui a permis de relever des défis de calcul considérée autrefois comme impossible.

Commençons par la description d'un scénario d'application qui nous accompagnera au long de cet article. Imaginons qu'il faille surveiller une chaîne de production (peu importe le produit). Nous pouvons acquérir des données provenant d'un large éventail de sources. Par exemple, nous pouvons placer des capteurs sur toute la chaîne de production, ou utiliser des *informations contextuelles*

qui indiquent l'âge et le type de chaque machine. Cet ensemble de données, ou jeu de *données*, peut être utilisé à différentes fins. Cela peut être la maintenance prédictive, pour évaluer et prévoir l'apparition de situations anormales, planifier les commandes de pièces de rechange ou entreprendre des réparations avant que les pannes ne se produisent, ce qui permet de réaliser des économies et d'accroître la productivité. En outre, la connaissance de l'historique des données nous permet de corrélérer les données mesurées par chaque capteur, pour mettre en évidence les éventuelles relations de cause à effet. Par exemple, si une augmentation soudaine de température et d'humidité dans la pièce a été suivie d'une diminution du nombre de pièces fabriquées, une modification pourrait consister à maintenir des conditions climatiques constantes à l'aide d'une climatisation.

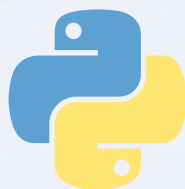
La mise en place d'un tel système n'est certainement pas à la portée de tous. Cependant, elle est simplifiée par les outils mis à disposition par la communauté *open source*. Il suffit d'avoir un PC (ou, à défaut, un Raspberry Pi si la quantité de données à traiter n'est pas énorme), de connaître Python (que vous pouvez approfondir en suivant un tutoriel comme celui-ci [2]) et, bien sûr les «outils» que je vous propose de découvrir ensemble !

Outils

Il faut savoir créer des programmes en Python. Pour cela, nous devons installer l'interpréteur que vous trouverez sur le site officiel de Python [3]. À partir d'ici, nous supposons que Python est installé et ajouté aux variables d'environnement de votre système.

L'environnement virtuel

Une fois l'installation de Python terminée, mettons en place un environnement virtuel, qui est en quelque sorte un «conteneur» séparé du reste de notre système et dans lequel sont installées les bibliothèques utilisées. La raison de l'utilisation d'un environnement virtuel pour l'installation globale des bibliothèques est liée à l'évolution rapide du monde Python. Très souvent, des différences substantielles surviennent même entre versions mineures de l'interpréteur, ce qui entraîne l'incompatibilité de bibliothèques



et, par conséquent, des programmes, car elles ont été écrites pour différentes versions de Python. En aménageant un environnement déterministe où la version de chaque bibliothèque installée est connue, nous aurons une sorte de garantie que nos programmes fonctionneront. En fait, il suffira de reproduire précisément la configuration de l'environnement virtuel et nous pouvons être sûrs que tout fonctionnera.

Pour gérer nos environnements virtuels, nous utilisons un logiciel appelé `virtualenvwrapper`. Celui-ci peut être installé à partir du shell en utilisant `pip` :

```
$ pip install virtualenvwrapper
```

Une fois l'installation terminée, nous créons un nouvel environnement virtuel comme suit :

```
$ mkvirtualenv ml-python
```

Notez que `ml-python` est le nom de l'environnement virtuel choisi pour notre exemple de scénario. Évidemment, ces noms peuvent varier et n'importe quel nom approprié pourra être choisi par le concepteur. Nous procédons ensuite à l'activation de l'environnement virtuel :

```
$ workon ml-python
```

Nous sommes prêts maintenant à installer les éléments nécessaires pour la suite de cet article.

Bibliothèques

Les bibliothèques présentées et utilisées ici sont les cinq plus utilisées pour l'analyse des données en Python.

La première, et peut-être la plus célèbre, est *NumPy*, qui est une sorte de port de MATLAB pour Python. *NumPy* est une bibliothèque pour calculs algébriques et matriciels. Par conséquent, les utilisateurs familiers de MATLAB trouveront de nombreuses similitudes, en termes de syntaxe et d'optimisation. L'utilisation du calcul algébrique dans *NumPy* est, en fait, plus efficace que les cycles imbriqués dans MATLAB (pour en savoir plus [4]). Comme on pouvait s'y attendre, le type de données au cœur de la fonction de *NumPy* est le *tableau* ou *array* en anglais. Il ne doit pas être

confondu avec le vecteur informatique correspondant, mais doit plutôt être compris au sens algébrique et géométrique du terme comme une *matrice*. Comme l'analyse des données est basée sur des opérations algébriques et matricielles, *NumPy* est également à la base de deux des cadres les plus utilisés : *scikit-learn* (dont il sera question plus loin) et *TensorFlow*.

Un complément naturel à *NumPy* est la bibliothèque *pandas*, qui gère et lit des données provenant de sources hétérogènes, y compris des tableaux Excel, des fichiers CSV, ou même des bases de données JSON et SQL. *pandas* est extrêmement flexible et puissant, vous permettant d'organiser les données en structures appelées *dataframes* manipulables selon les besoins et exportables facilement directement dans des tableaux *NumPy*.

La troisième bibliothèque que nous utiliserons est *scikit-learn*. Née d'un projet universitaire, *scikit-learn* est un cadre qui met en œuvre la plupart des algorithmes d'apprentissage automatique utilisés de nos jours en fournissant une interface commune. Ce dernier concept est précisément celui de la programmation orientée objet : il est en effet possible d'utiliser pratiquement tous les algorithmes proposés par *scikit-learn* grâce à la méthode `fit_transform` en passant au moins deux paramètres, tels que les données analysées et les étiquettes associées.

Les deux dernières bibliothèques que nous utiliserons sont *Matplotlib* et *Jupyter*. La première, avec son complément *Seaborn*, est nécessaire pour visualiser les résultats de nos expériences sous forme de graphiques. La seconde nous offre l'utilisation de *notebooks* ou *carnets*, environnements interactifs d'utilisation simple et immédiate qui permettent à l'analyste de données d'écrire et d'exécuter des parties de code indépendamment des autres.

Avant d'aller plus loin, cependant, voici quelques concepts théoriques nécessaires à la construction d'une «base commune» pour le discours.

Les concepts

Le premier concept requis est celui des *ensembles de données* (ou *datasets*), ce qui est souvent considéré comme allant de soi. Il s'agit d'ensembles d'échantillons, chacun d'entre eux étant caractérisé par



Figure 1. L'écran d'accueil pour la gestion des carnets (notebook) dans Jupyter.

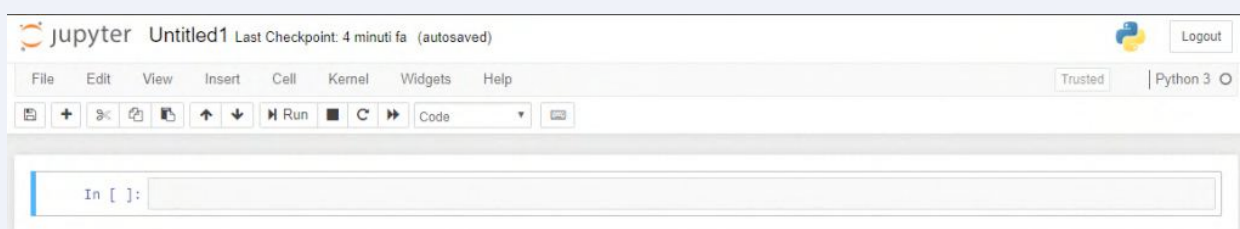


Figure 2. Un carnet (notebook) vide.

un certain nombre de variables ou de caractéristiques, qui décrivent le phénomène observé. Par souci de simplicité, nous pouvons considérer un ensemble de données comme un tableau Excel. Les lignes fournissent les *échantillons*, c'est-à-dire les *observations individuelles* du phénomène, tandis que les colonnes fournissent les *caractéristiques*, c'est-à-dire les valeurs qui caractérisent chacun des aspects du processus. Pour revenir à l'exemple de la fabrication intelligente, chaque ligne représentera les conditions de la chaîne de production à un moment donné tandis que chaque colonne comportera la valeur lue sur un capteur donné.

À propos de scikit-learn, il y a lieu de mentionner le concept de *label* ou de *classe*. La présence ou l'absence d'étiquettes permet de distinguer les algorithmes supervisés des algorithmes non supervisés. La différence est, du moins en principe, assez simple : les *algorithmes supervisés* requièrent une connaissance a priori de la classe de chaque échantillon de l'ensemble de données de l'exemple, pas les *algorithmes non supervisés*. En pratique, pour utiliser un algorithme supervisé, il est nécessaire qu'un expert du domaine établisse la *classe d'appartenance* de chaque échantillon. Dans le cas d'un processus de fabrication intelligent, un «expert» pourrait déterminer si un ensemble de lectures, à un moment précis, représente une situation anormale ou non. Ainsi, l'échantillon unique peut être associé à l'une de ces deux classes possibles (anormal/normal). Cela n'est pas nécessaire pour les algorithmes non supervisés.

Il faut aussi distinguer les processus avec des données *indépendantes et identiquement distribuées* (IID) des données dans un *ordre chronologique*. La différence est liée à la nature du phénomène observé. Les échantillons d'un processus IID sont indépendants les uns des autres, alors que dans une série temporelle, chaque échantillon dépend d'une combinaison linéaire ou non linéaire des valeurs que le processus a produites à des moments précédents.

Commençons !

Maintenant que les termes théoriques et pratiques requis sont connus, nous passons à l'utilisation d'un ensemble de données approprié pour notre exemple. L'ensemble de données utilisé est SECOM, un acronyme qui signifie SEmiCOnductor Manufacturing, qui contient les valeurs lues par un ensemble de capteurs lors de la

surveillance d'un processus de fabrication de semi-conducteurs. Dans l'ensemble de données, téléchargeable à partir de différentes sources (comme Kaggle [5]), il y a 590 variables, représentatives chacune de la lecture d'un seul capteur à un moment donné. L'ensemble de données contient également des étiquettes qui permettent de distinguer les défaillances et les anomalies du bon fonctionnement du système.

Une fois l'ensemble de données téléchargé, nous installons les bibliothèques mentionnées ci-dessus. Depuis la ligne de commande, saisissez :

```
$ pip install numpy pandas scikit-learn matplotlib
seaborn jupyter numpy pandas install
```

Une fois les bibliothèques installées, nous pouvons mettre en place un simple pipeline pour l'analyse des données.

Le premier carnet

La première étape consiste à créer un nouveau *carnet*. À partir de la ligne de commande, nous lançons Jupyter par les instructions suivantes :

```
$ jupyter-notebook
```

Un écran similaire à celui de la **figure 1** s'ouvre. Nous créons un *carnet* en sélectionnant *Nouveau > Python 3*. Un nouvel onglet s'ouvrira dans notre navigateur avec le carnet créé. Prenons le temps de nous familiariser avec l'interface (**fig. 2**) qui ressemble (très vaguement) à une ligne de commande interactive, un menu et plusieurs options.

Ce qui saute aux yeux est la *cellule*. L'exécution de cellules individuelles est lancée par le bouton *Run*. Elle est indépendante de celle des autres cellules (gardons à l'esprit la validité du concept de *portée des variables*).

Les trois boutons immédiatement à droite du bouton *Run* permettent d'arrêter, de redémarrer et d'initialiser le *noyau*, c'est-à-dire l'instance associée à notre carnet par Jupyter. Le redémarrage de l'instance peut être nécessaire pour réinitialiser les variables locales et globales associées au script, ce qui est particulièrement utile lorsque vous expérimentez avec de nouvelles méthodes et bibliothèques.

	a21	a87	a88	a89	a114	a115	a116	a117	a118	a120	...	a528
count	1567.000000	1567.000000	1567.000000	1567.000000	1567.000000	1567.000000	1567.000000	1567.000000	1567.000000	1567.000000	...	1567.000000
mean	1.405054	2.401872	0.982420	1807.815021	0.945424	0.000123	747.383792	0.987130	58.625908	0.970777	...	6.395717
std	0.016737	0.037332	0.012848	53.537262	0.012133	0.001668	48.949250	0.009497	6.485174	0.008949	...	1.888698
min	1.179700	2.242500	0.774900	1627.471400	0.853400	0.000000	544.025400	0.890000	52.806800	0.841100	...	2.170000
25%	1.396500	2.376850	0.975800	1777.470300	0.938600	0.000000	721.023000	0.989500	57.978300	0.964800	...	4.895450
50%	1.406000	2.403900	0.987400	1809.249200	0.946400	0.000000	750.861400	0.990500	58.549100	0.969400	...	6.410800
75%	1.415000	2.428600	0.989700	1841.873000	0.952300	0.000000	776.781850	0.990900	59.133900	0.978300	...	7.594250
max	1.453400	2.555500	0.993500	2105.182300	0.976300	0.041400	924.531800	0.992400	311.734400	0.982700	...	14.447900

Figure 3. Les cinq premières lignes de l'ensemble de données SECOM.

	a1	a2	a3	a4	a5	a6	a7	a8	a9	a10	...	a582	a583	a584	a585	a586	a587	a588	a589
0	3030.93	2564	2187.7333	1411.1265	1.3602	100	97.6133	0.1242	1.5005	0.0162	...	?	0.5005	0.0118	0.0035	2.363	?	?	?
1	3095.78	2465.14	2230.4222	1463.6606	0.8294	100	102.3433	0.1247	1.4966	-0.0005	...	208.2045	0.5019	0.0223	0.0055	4.4447	0.0096	0.0201	0.004
2	2932.61	2559.94	2186.4111	1698.0172	1.5102	100	95.4878	0.1241	1.4436	0.0041	...	82.8602	0.4958	0.0157	0.0039	3.1745	0.0584	0.0484	0.014
3	2988.72	2479.9	2199.0333	909.7926	1.3204	100	104.2367	0.1217	1.4882	-0.0124	...	73.8432	0.499	0.0103	0.0025	2.0544	0.0202	0.0149	0.004
4	3032.24	2502.87	2233.3667	1326.52	1.5334	100	100.3967	0.1235	1.5031	-0.0031	...	?	0.48	0.4766	0.1045	99.3032	0.0202	0.0149	0.004

5 rows x 591 columns

Figure 4. Brève description statistique de l'ensemble de données SECOM.

Une autre option utile est le choix possible du type de cellule, entre *Code* (c'est-à-dire code Python), *Markdown* (utile pour insérer des commentaires et des descriptions dans le format utilisé, par exemple, par GitHub READMEs), *Raw NBContent* (texte brut) et *Heading* (offrant un raccourci pour insérer des titres).

Importation et affichage des données

Une fois familiarisés avec l'interface, nous pouvons passer à la mise en œuvre de notre script. Ici, nous importons les bibliothèques et les modules que nous utiliserons :

```
import numpy as np
import pandas as pd
import matplotlib.pyplot as plt
%matplotlib inline
import seaborn as sns
from ipywidgets import interact
from sklearn.ensemble import RandomForestClassifier
from sklearn.impute import SimpleImputer
from sklearn.model_selection import train_test_split
from sklearn.preprocessing import StandardScaler
from sklearn.metrics import confusion_matrix,
accuracy_score
from sklearn.utils import resample
```

Signalons l'instruction `%matplotlib inline` qui nous permet d'afficher correctement les graphiques produits par Matplotlib. Ensuite, les données du fichier contenant l'ensemble de données SECOM sont importées à l'aide de la fonction `read_csv` de *pandas*. Notez que, dans cet exemple, pour simplifier, le chemin relatif du fichier est codé en dur. Cependant, il est conseillé d'utiliser le paquet *os* Python pour permettre à notre programme de déterminer lui-même ce chemin lorsque c'est nécessaire.

```
data = pd.read_csv('data/secom.csv')
```

L'instruction précédente lit les données contenues dans le fichier *secom.csv*, en les organisant dans un cadre de données nommé *data*. Nous pouvons afficher les cinq premières lignes de la trame de données grâce à l'instruction `head()` (fig. 3).

```
data.head()
```

La visualisation des premières lignes du cadre de données peut être utile pour obtenir un premier aperçu des données à analyser. Dans ce cas, on remarque immédiatement la présence de certaines valeurs égales à « ? » qui représentent probablement des valeurs nulles. En outre, il est évident que la plage des valeurs varie considérablement, un facteur que nous devrons garder à l'esprit plus tard. Nous pouvons également utiliser la fonction `describe()` pour obtenir un aperçu rapide des caractéristiques statistiques de chaque variable (fig. 4).

```
data.describe()
```

L'analyse statistique peut, en général, mettre en évidence des conditions présentant un manque de normalité (c'est-à-dire des données dont la distribution n'est pas paramétrique), ou la présence d'anomalies. À titre d'exemple, nous constatons que l'écart-type (std) associé aux variables *a116* et *a118* est assez élevé proportionnellement, de sorte que nous nous attendons à une forte signification de ces variables lors de leur analyse. D'autre part, des variables telles que *a114* ont un faible std, on s'attend donc à ce qu'elles soient écartées faute de pertinence dans le processus analysé.

Une fois le chargement et l'affichage de la trame de données terminés, nous pouvons passer à une partie fondamentale du pipeline : le *prétraitement*.

Prétraitement des données

Dans un premier temps, nous affichons le nombre d'échantillons associés à chaque classe. Pour ce faire, nous utiliserons la fonction `value_counts()` sur la colonne `classvalue` car elle contient les étiquettes associées à chaque échantillon.

```
data[«classvalue»].value_counts()
```

Nous voyons qu'il y a 1463 échantillons prélevés en situation normale de fonctionnement (*classe -1*) et 104 en situation de défaillance (*classe 1*). L'ensemble de données est donc fortement déséquilibré et il serait approprié de prendre des mesures pour uniformiser la répartition des échantillons entre différentes classes. Cela renvoie à la fonction intrinsèque des algorithmes d'apprentissage machine qui apprennent sur la base des données dont ils

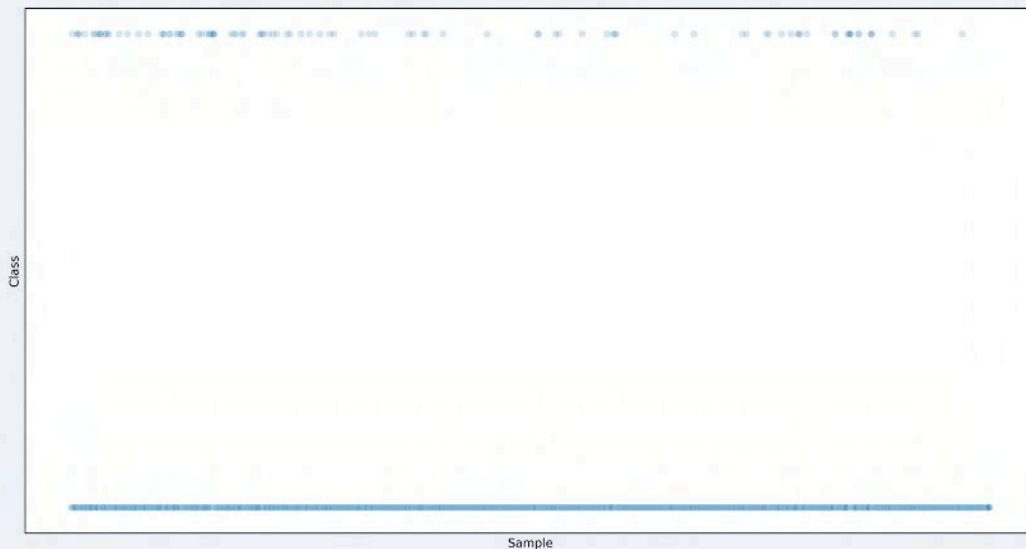


Figure 5. Nombre d'échantillons par classe dans l'ensemble de données SECOM.

disposent. Dans ce cas précis, l'algorithme apprendra à caractériser avec succès une situation de comportement standard, mais sa caractérisation des situations anormales sera affectée par de «incertitudes». Le déséquilibre est encore plus évident lorsque l'on observe le nuage de points (**fig. 5**) :

```
sns.scatterplot(data.index, data['classvalue'],
               alpha=0.2)
plt.show()
```

En gardant à l'esprit ce déséquilibre (auquel nous reviendrons), nous séparons les étiquettes des données :

```
labels = data['classvalue']
data.drop('classvalue', axis=1, inplace=True)
```

Notez l'utilisation du paramètre `axis` dans la fonction `drop` qui nous permet de spécifier que la fonction doit opérer sur les colonnes du cadre de données (`dataframe`). Par défaut, les fonctions `pandas` opèrent sur les lignes.

Un autre aspect qui peut être extrapolé de l'analyse des ensembles de données est la présence, dans cette version spécifique des données SECOM, de nombreuses colonnes contenant des données de différents types, c'est-à-dire à la fois des chaînes de caractères et des chiffres. Par conséquent, les `pandas`, incapables de déterminer avec certitude le type de données avec lequel chaque caractéristique est représentée, en reportent la définition sur l'utilisateur. Pour mettre toutes les données sous forme numérique, il faut donc utiliser trois fonctions offertes par les `pandas`.

La première fonction que nous utiliserons est `replace()`, avec laquelle nous pouvons remplacer tous les points d'interrogation par la valeur constante `numpy.nan`, le caractère de remplacement utilisé pour traiter les valeurs nulles dans les `tableaux` NumPy.

```
data = data.replace('?', np.nan, regex=False)
```

Le premier paramètre de la fonction est la valeur à remplacer, le deuxième est la valeur à utiliser pour le remplacement, et le troisième est un drapeau indiquant si le premier paramètre représente ou non une expression régulière. Nous pourrions également utiliser une syntaxe alternative en utilisant le paramètre `inplace` réglé sur `True`, comme suit :

```
data.replace('?', np.nan, regex=False, inplace=True)
```

Les deuxième et troisième fonctions que nous pouvons utiliser pour résoudre les problèmes signalés ci-dessus sont les fonctions `apply()` et `to_numeric()` respectivement. La première permet d'appliquer une certaine fonction à toutes les colonnes (ou lignes) d'une trame de données, tandis que la deuxième convertit une seule colonne en valeurs numériques. En les combinant, nous produisons des données uniques et nous supprimons également les valeurs que NumPy et scikit-learn ne peuvent traiter :

```
data.apply(pd.to_numeric)
```

Nous devons maintenant évaluer quelles caractéristiques contenues dans l'ensemble de données sont réellement utiles. Nous utilisons généralement des techniques (plus ou moins complexes) de *sélection des caractéristiques* pour réduire les redondances et la taille du problème à traiter, ce qui présente des avantages évidents en termes de temps de traitement et de performance de l'algorithme. Dans notre cas, nous nous appuyons sur une technique moins complexe qui implique l'élimination des caractéristiques de variance faible (et donc, comme mentionné, de faible importance). Nous créons donc un widget interactif qui nous permet de visualiser, sous la forme d'un histogramme, la distribution des données pour chaque caractéristique :

```
@interact(col=(0, len(df.columns) - 1))
def show_hist(col=1):
    data['a' + str(col)].value_counts().
    hist(grid=False, figsize=(8, 6))
```

L'interactivité est assurée par le décorateur (= *decorator*) `@interact`, dont la valeur de référence (c'est-à-dire `col`) varie entre 0 et le nombre de caractéristiques présentes dans l'ensemble de données. En explorant les données affichées par le widget, nous déterminerons combien de caractéristiques prennent une valeur unique, ce qui signifie qu'elles peuvent être simplement ignorées dans l'analyse. Nous pouvons alors décider de les éliminer comme suit :

```
single_val_cols = data.columns [len(data)/data.
                               nunique() < 2]
secom = data.drop(single_val_cols, axis=1)
```

Bien sûr, il existe des techniques de sélection de caractéristiques plus pertinentes et plus raffinées qui utilisent par exemple des paramètres statistiques. La documentation scikit-learn [6] en donne un aperçu complet.

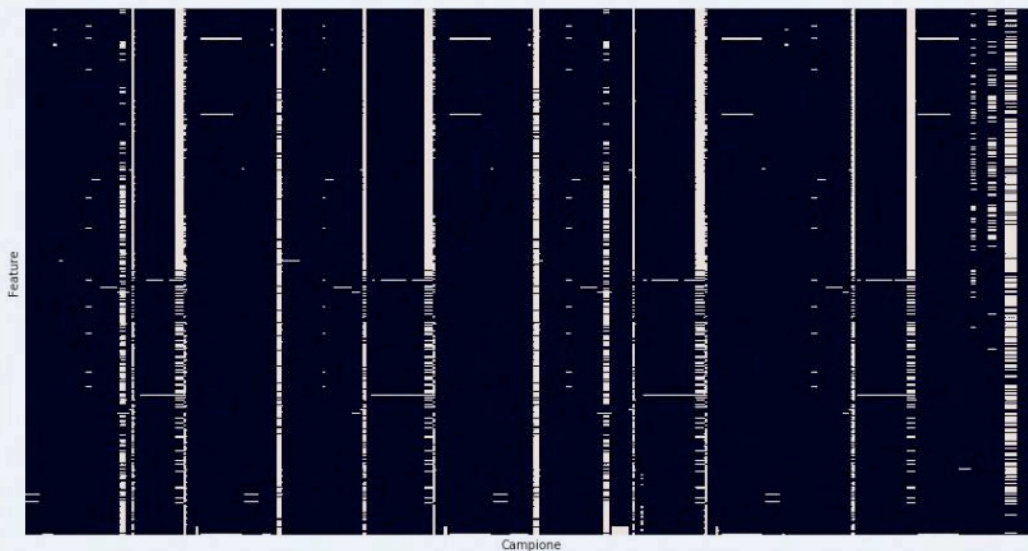


Figure 6. Carte des valeurs nulles dans l'ensemble de données SECOM.

La dernière étape consiste à traiter les valeurs nulles (que nous avons remplacées précédemment par `np.nan`). Nous inspectons l'ensemble de données pour voir combien il y en a ; pour ce faire, nous utilisons une carte thermique (fig. 6) où les points blancs représentent les valeurs nulles.

```
sns.heatmap(secom.isnull(), cbar=False)
```

Il est évident que de nombreux échantillons présentent un pourcentage élevé de valeurs nulles qui ne doivent pas être prises en compte afin d'éliminer l'effet de *biais* sur les données.

```
na_cols = [col for col in secom.columns if
            secom[col].isnull().sum() / len(secom) > 0.4]
secom_clean = secom.drop(na_cols, axis=1)
secom_clean.head()
```

Grâce aux commandes précédentes, une *liste de compréhension* (= *comprehension list*) a maintenant isolé toutes les caractéristiques ayant plus de 40 % de valeurs nulles, ce qui permet de les supprimer de l'ensemble de données.

Les caractéristiques ayant moins de 40 % de valeurs nulles doivent encore être traitées. Nous pouvons utiliser ici notre premier objet scikit-learn, le `SimpleImputer`, qui attribue des valeurs à tous les NaNs en fonction d'une stratégie définie par l'utilisateur. Dans ce cas, nous utiliserons une stratégie *moyenne*, associant la valeur moyenne supposée par l'objet à chaque NaN.

```
imputer = SimpleImputer(strategy='mean')
secom_imputed = pd.DataFrame(imputer.
                              fit_transform(secom_clean))
secom_imputed.columns = secom_clean.columns
```

À titre d'exercice, nous vérifions que nous n'avons aucun zéro dans l'ensemble de données avec une autre carte thermique qui, comme on peut s'y attendre, sera uniformément foncée. Ensuite, nous pouvons passer au traitement proprement dit.

Traitement des données

Divisons notre ensemble de données en deux sous-ensembles : un ensemble de *formation* et un ensemble de *test*. Cette subdivision est nécessaire pour atténuer le phénomène de *surapprentissage* (= *overequipment*), par lequel l'algorithme adhère trop aux données [7], et garantit l'applicabilité du modèle à des cas différents de ceux pour lesquels il a été formé. Pour ce faire, nous utilisons

la fonction `train_test_split` :

```
X_train, X_test, y_train, y_test = train_test_
split(secom_imputed, labels, test_size=0.3)
```

Le paramètre `test_size` permet de spécifier le pourcentage de données réservées au test ; les valeurs standard de ce paramètre se situent généralement entre 0,2 et 0,3.

Il est également important de *normaliser* les données. Nous avons remarqué que les valeurs de certaines caractéristiques présentent des variations beaucoup plus fortes que d'autres, ce qui leur donne plus de poids. Le fait de les normaliser les ramène dans une seule plage de valeurs afin d'éviter les déséquilibres dus aux décalages initiaux. Pour ce faire, nous utilisons `StandardScaler` :

```
scaler = StandardScaler()
X_train = pd.DataFrame(scaler.fit_transform(X_train),
                        index=X_train.index, columns=X_train.columns)
X_test = pd.DataFrame(scaler.fit_transform(X_test),
                       index=X_test.index, columns=X_test.columns)
```

Il est intéressant de noter l'utilité de l'interface commune offerte par scikit-learn. Tant le `scaler` que l'`imputer` utilisent la méthode `fit_transform` pour traiter les données qui, dans des pipelines complexes, simplifient grandement l'écriture du code et la compréhension de la bibliothèque.

Nous voici enfin prêts à classer les données. En particulier, nous utiliserons une *forêt aléatoire* (*random forest*) [8], en obtenant, après formation, un modèle capable de distinguer les situations normales et anormales. Nous allons vérifier les performances du modèle identifié de deux manières. La première est le *score de précision* (*accuracy score*). Il s'agit du pourcentage d'échantillons appartenant à la série de tests correctement classés par l'algorithme. La seconde est la *matrice de confusion* [9] qui met en évidence le nombre de faux positifs et de faux négatifs.

Créons d'abord le classificateur :

```
clf = RandomForestClassifier(n_estimators=500,
                             max_depth=4)
```

Cela crée une forêt aléatoire avec 500 estimateurs dont la profondeur maximale est de 4 niveaux. Nous pouvons maintenant entraîner notre modèle sur des données d'entraînement :

```
clf.fit(X_train, y_train)
```

Une fois la formation terminée, le modèle formé est utilisé pour classer les échantillons d'essai :

```
y_pred = clf.predict(X_test)
```

Il en résulte deux étiquettes qui se rapportent à chacun des deux tests. Le premier, qui appartient à `y_test`, représente la «vérité», tandis que le second, qui appartient à `y_pred`, est la valeur prédite par l'algorithme. En les comparant, nous déterminons à la fois `accuracy` et `confusion_matrix`.

```
accuracy = accuracy_score(y_test, y_pred)
cf = confusion_matrix(y_test, y_pred)
```

L'exemple donne les résultats suivants :

```
Model accuracy on test set is: 0.9341825902335457
The confusion matrix of the model is:
[[440   0]
 [ 31   0]]
```

La précision, d'environ 93 %, est excellente, donc le modèle semble bon. Cependant, nous constatons que le modèle est très précis pour la classification des échantillons de la classe prédominante, mais très imprécis pour la classification des échantillons de la classe minoritaire. Donc il y a manifestement un biais. Il nous faut une stratégie pour améliorer cette situation. C'est possible en suréchantillonnant les données appartenant à la classe minoritaire de manière à équilibrer, au moins partiellement, l'ensemble des données. Pour ce faire, nous utiliserons la fonction `resample` (rééchantillonnage) de `pandas`.

```
normals = data[data['classvalue'] == -1]
anomalies = data[data['classvalue'] == 1]
anomalies_upsampled = resample(anomalies,
                               replace=True, n_samples=len(normals))
```

Cela permet d'augmenter la taille de l'ensemble de données afin de rapprocher le nombre d'échantillons normaux du nombre d'échantillons anormaux. En conséquence, nous devons redéfinir `X` et `Y` comme suit :

```
upsampled = pd.concat([normals, anomalies_upsampled])
X_upsampled = upsampled.drop('classvalue', axis=1)
y_upsampled = upsampled['classvalue']
```

En effectuant à nouveau la formation (y compris en répétant les procédures de fractionnement et de normalisation), nous obtenons les résultats suivants pour notre exemple :

```
Model accuracy on test set is: 0.8631921824104235
The confusion matrix of the model is:
[[276  41]
 [ 43 254]]
```

Visiblement la précision du modèle a diminué, probablement à cause de la plus grande hétérogénéité induite dans l'ensemble des données. Cependant, en regardant la matrice de confusion, nous remarquons immédiatement que le modèle a en réalité amélioré ses capacités de généralisation, réussissant également à classer correctement les échantillons appartenant à des situations anormales.

Conclusions et références

Dans cet article, nous avons présenté un pipeline pour l'analyse de données provenant de processus réels en Python. Manifestement chacun des sujets abordés est extrêmement vaste. Une expérience théorique et pratique est essentielle si vous voulez vous engager sérieusement dans l'analyse de données. Nous avons appris également qu'il ne faut pas s'arrêter au premier résultat obtenu, même dans des situations aussi complexes que celle dont il est question. Il convient d'interpréter de différents points de vue les résultats obtenus afin de distinguer un modèle opérationnel d'un modèle plus ou moins évidemment faussé.

Le message à retenir est donc le suivant : l'analyse des données ne peut être une discipline mécanique, elle exige au contraire une analyse critique, approfondie et variée du phénomène observé, guidée par des notions théoriques et des compétences pratiques. Je recommande vivement les références ci-dessous grâce auxquelles vous pourrez approfondir vos connaissances sur certains des aspects abordés dans l'article, ainsi que le lien vers le dépôt du GitLab où vous pourrez consulter le code écrit pour cet article. 

200505-03

Article publié en italien par Elettronica Open Source (<https://it.emcelettronica.com>) et traduit par Elektor avec sa permission.

LIENS

- [1] **Support informatique GPU MATLAB** : <https://uk.mathworks.com/solutions/gpu-computing.html>
- [2] **Guides pour les débutants en programmation Python** : <https://wiki.python.org/moin/BeginnersGuide/Programmers>
- [3] **Python** : www.python.org/
- [4] **Bonnes pratiques d'optimisation dans MATLAB** : <https://uk.mathworks.com/videos/best-practices-for-optimisation-in-matlab-96756.html>
- [5] **Ensemble de données UCI SECOM** : <https://www.kaggle.com/paresh2047/uci-semcom/kernels>
- [6] **Guide de l'utilisateur Scikit-Learn** : https://scikit-learn.org/stable/user_guide.html
- [7] **Sur- ou sous-apprentissage – un exemple complet** : <https://towardsdatascience.com/overfitting-vs-underfitting-a-complete-example-d05dd7e19765>
- [8] **Comprendre la forêt aléatoire** : <https://towardsdatascience.com/understanding-random-forest-58381e0602d2>
- [9] **Comprendre la matrice de confusion** : <https://towardsdatascience.com/understanding-confusion-matrix-a9ad42dcfd62>
- [10] **Dépôt GitLab pour cet article** : <https://gitlab.com/eos-acard/machine-learning-in-python>

Rémy Mallard explique les oscilloscopes anciens & modernes pour les débutants

Points forts

- Livre orienté vers la pratique, avec « juste ce qu'il faut » de théorie
- Traite aussi le sujet des oscilloscopes numériques modernes
- Ouvrage utile avant l'achat pour guider le choix d'un oscilloscope
- Manuel d'utilisation à consulter régulièrement après l'achat
- Comme dans ses précédents livres, l'auteur partage sa bonne humeur communicative et sa grande expérience



elektor

En électronique, si l'on veut progresser dans le plaisir et dans la compréhension, il faut un oscillo. Comment le choisir ? À peine cette question-là aura-t-elle trouvé sa réponse, il en viendra d'autres qui se résument ainsi : comment se servir de l'oscilloscope de telle sorte que ce qu'il affiche corresponde à la réalité des signaux ? Dans ce livre, Rémy Mallard, répond à ces questions.

Auteur : Rémy Mallard
Pages : 375 (en couleur)
Format : 17 x 23,5 cm (broché)
ISBN : 978-2-86661-208-5

Pour commander ce livre :
www.elektor.fr/19124

