

Analyse Discriminante Linéaire (LDA) classique afin de trouver la meilleure projection dans l'espace de dimension inférieure avant de projeter sur l'hyper-sphère. Le but de ce travail est de proposer une projection non linéaire directement de l'espace des i -vectors vers une hypersphère en minimisant la corrélation inter-classe tout en maximisant la corrélation intra-classe. Il sera question de comparer les résultats obtenus avec d'autres solutions comme : le CD classique, le PLDA et l'approximation de mesure de distance avec le perceptron multicouche et les machines de Boltzmann.

Climb

La programmation orientée contexte appliquée au traitement d'images

par François Ripault © 11 juillet, 11h00, Amphi 3

La programmation orientée contexte est un pa-

radigme permettant de prendre en compte les problématiques transverses et les variations de comportement d'un programme dépendantes du contexte. Ce paradigme permet ainsi d'exprimer des aspects du comportement d'un programme qui sont orthogonaux à sa modélisation objet, tout en maintenant son abstraction et sa modularité. Dans le domaine du traitement d'images, par exemple, la programmation orientée contexte peut être utilisée pour prendre en compte des aspects qui traitent à la fois de la structure de l'image, de son contenu, ou encore sa représentation mémoire.

Nous présentons la programmation orientée contexte, puis nous analysons ensuite les problématiques transverses présentes dans Climb, une bibliothèque de traitement d'images écrite en Common Lisp. Nous expliquons ensuite comment la programmation orientée contexte permet de résoudre ces problématiques. Enfin, nous analysons les avantages de la programmation orientée contexte dans le domaine du traitement d'images.

Contacter le LRDE



18, rue Pasteur
Paritalie, bâtiment X, aile Mistral
2e étage, droite droite

Tél. : 01 53 14 59 22

Fax : 01 53 14 59 13

Contact : info@lrde.epita.fr

Les permanents : lrde@lrde.epita.fr

Site Web : <http://www.lrde.epita.fr>

... et surtout, passez nous voir ;
vous serez toujours les bienvenus !



L'aléastriel du Laboratoire de Recherche et de Développement de l'EPITA¹

Numéro 28, Juin 2013

Édito

par Raphaël Boissel (CSI 2014)

Ce numéro spécial "Séminaires CSI" a pour objectif de communiquer sur les horaires et les sujets de toutes les présentations qui seront faites par les étudiants. En effet, comme chaque année, nous (les étudiants du LRDE) présentons nos travaux de recherche réalisés au cours de ce semestre. Les présentations seront réparties sur deux journées. L'une sera dédiée au séminaire des élèves de première an-

L'air de rien N° 28 Séminaires CSI de juillet 2013

née qui aura lieu le 3 juillet et l'autre sera dédiée aux élèves de deuxième année le 11 juillet.

Nous espérons que vous viendrez nombreux à ces présentations qui représentent non seulement pour nous un moyen de partager avec vous les résultats de nos travaux, mais, qui constituent aussi pour les étudiants toutes années confondues une expérience enrichissante et formatrice.

Programme des séminaires CSI

Séminaire du 3 juillet

3 juillet 09h30–13h15 — amphi 3 KB	
Olena	
9h30	Analyse structurelle haut-niveau de document dans le module Scribo d'Olena — <i>Christophe Escobar</i>
10h00	Détection de logotype et autres invariants caractéristiques — <i>Anthony Seure</i>
Speaker ID	
10h30	Les Machines de Boltzmann dans la reconnaissance du locuteur — <i>Jean-Luc Bounthing</i>
11h00	Approximation de la distance entre i -vecteurs par Perceptron Multi-Couches. — <i>Jimmy Yeh</i>
Climb	
11h45	Maintenance automatique des symboles exportés dans les packages de Common Lisp. — <i>Christophe Vermorel</i>
Vaucanson	
12h15	Améliorer l'architecture de Vaucanson 2 — <i>Guillaume Sanchez</i>
12h45	Transducteurs dans Vaucanson 2 — <i>Victor Santet</i>

Séminaire du 11 juillet

11 juillet 09h30–11h30 — amphi 3 KB	
Olena	
9h30	Réduire les ressources utilisées par une chaîne de traitement complexe — <i>Raphaël Boissel</i>
10h00	À propos du calcul de l'arbre des formes sur des images n -dimensionnelles en temps quasi-linéaire — <i>Sébastien Crozet</i>
Speaker ID	
10h30	Projection non-linéaire pour l'attribution de score selon la distance en cosinus dans le contexte des systèmes de vérification du locuteur à base d' i -vecteurs — <i>Benjamin Roux</i>
Climb	
11h00	La programmation orientée contexte appliquée au traitement d'images — <i>François Ripault</i>

1. L'air de rien, <http://publis.lrde.epita.fr/LrdeBulletin>.

Résumés

Séminaire du 3 Juillet

Olena

Analyse structurelle haut-niveau de document dans le module Scribo d'Olena

par *Christophe Escobar* 3 juillet, 09h30, Amphi 3

L'extraction de structures dans un document numérisé se base sur la mise en place d'une chaîne de traitements constituée de briques élémentaires. L'analyse haut-niveau d'un document nécessite des informations structurelles sur celui-ci et se basera donc sur cette chaîne de traitements. Elle consistera à extraire des informations plus abstraites de nature structurelle sur un document, pour obtenir des "indices" sur sa structure. À l'aide de ces indices et de schémas de structure, il est ensuite possible de réaliser des traitements de haut-niveau tels que l'identification du flot de lecture, l'extraction d'éléments spécifiques ou la reconstruction d'un document dans un autre format.

Détection de logotypes et autres invariants caractéristiques

par *Anthony Seure* 3 juillet, 10h00, Amphi 3

La détection de logotypes et d'invariants dans une image a pour but de trouver parmi une ou plusieurs images un élément graphique qui caractérise une marque, entreprise, etc. De tels éléments peuvent se retrouver dans de nombreuses images naturelles mais également dans des images publicitaires. L'intérêt de l'intégration d'un tel outil dans Olena et Terra Rush permettrait de mieux indexer les contenus mais également d'invalider des zones de l'image dans d'autres chaînes de traitement. Nous expliquerons principalement une méthode générique permettant de localiser les points-clés invariants d'une image : les descripteurs SIFT.

Speaker ID

Les Machines de Boltzmann dans la reconnaissance du locuteur

par *Jean-Luc Bounthong* 3 juillet, 10h30, Amphi 3

L'espace Total Variability (TV) représente actuellement l'état de l'art dans le domaine de la vérification du locuteur. Des progrès significatifs ont été réalisés grâce aux nouvelles méthodes de classification comme l'Analyse discriminante linéaire proba-

biliste (PLDA) ou la Distance Cosinus (CD). Dans cette étude, nous explorons une nouvelle méthode pour déterminer la distance entre deux i-vecteurs en utilisant une variante des Machines de Boltzmann, ces nouvelles approches ont montré des résultats satisfaisants dans le domaine du traitement d'images. Nous allons aussi comparer les performances en terme de fiabilité avec les méthodes classiques comme PLDA ou CD.

Approximation de la distance entre i-vecteurs par Perceptron Multi-Couches

par *Jimmy Yeh* 3 juillet, 11h00, Amphi 3

Actuellement, l'espace des i-vecteurs est la représentation standard des paramètres de la parole dans les systèmes de reconnaissance du locuteur. Le calcul du score est généralement basé sur la distance cosinus, ou sur l'analyse discriminante linéaire probabiliste. Le but de ce sujet est de remplacer ces approches par un Perceptron Multi-Couches (PMC). Le PMC a montré en effet de bonnes performances pour approximer des fonctions non linéaires. L'idée principale étant de trouver une meilleure fonction que la distance cosinus. Les performances du PMC seront comparées aux autres méthodes comme la distance cosinus ou encore la machine de Boltzmann.

Climb

Maintenance automatique des symboles exportés dans les packages de Common Lisp

par *Christophe Vermorel* 3 juillet, 11h45, Amphi 3

Les "packages" de Common Lisp offrent une fonctionnalité analogue aux espaces de noms présents dans des langages comme le C++. Ceux-ci permettent d'encapsuler des symboles qui peuvent être exportés ou privés. Les symboles exportés peuvent être explicitement déclarés lors de la définition du package. Cette liste de symboles est fastidieuse à maintenir lors du développement de projets de grosse envergure. Dans ce rapport, nous étudions des techniques de maintenance automatiques de cette liste. Plusieurs alternatives sont présentées et comparées.

Vaucanson

Améliorer l'architecture de Vaucanson 2

par *Guillaume Sanchez* 3 juillet, 12h15, Amphi 3

Vaucanson 2 est la suite de Vaucanson, la plateforme de manipulation d'automates finis. Ce redémarrage à zéro a été fortement encouragé par quelques problèmes de conception. Vaucanson 2 entend apprendre des erreurs de son prédécesseur, et dans ce but, un vrai travail de conception doit être fait. Les templates doivent être utilisés pour la performance (évitant un usage abusif de fonctions virtuelles) afin de pouvoir préférer la généralité plutôt que la généralité, tout en conservant une flexibilité lors de l'exécution. Vaucanson 2 prétendra aussi à une compilation dynamique de code (types d'automates et algorithmes) et de le charger en tant que bibliothèque dynamique. Ces simples problèmes amènent à des architectures de distribution complexes et d'effacement de type, qui doivent être testés dans un environnement simulé.

Transducteurs dans Vaucanson 2

par *Victor Santet* 3 juillet, 12h45, Amphi 3

Vaucanson est une bibliothèque dédiée à la manipulation d'automates finis développée par le Laboratoire de Recherche et Développement de l'EPITA (LRDE). Ce projet tient à offrir à ses utilisateurs une large variété de types d'automates, notamment les transducteurs. Ce rapport propose une implémentation de transducteurs génériques, qui peuvent accepter des automates étiquetés par des n -uplets de langages.

Séminaire du 11 Juillet

Olena

Réduire les ressources utilisées par une chaîne de traitement complexe

par *Raphaël Boissel* 11 juillet, 09h30, Amphi 3

Même si la précision du système de localisation de texte développé par le laboratoire a aujourd'hui grandement augmentée, un des principaux problèmes demeure la vitesse. En effet, si l'on souhaite être en mesure de l'utiliser pour extraire du texte dans des vidéos ou s'il doit être embarqué dans des appareils aux ressources limitées tel que des téléphones portables, la consommation de ressources mémoire et CPU doit être réduite.

Pour parvenir à cet objectif, plusieurs méthodes permettant de réduire l'impact mémoire et CPU des étapes les plus coûteuses peuvent être utilisées. Cependant, toutes n'offrent pas les mêmes résultats en fonction de l'image d'entrée et certaines peuvent même dégrader le résultat final produit par le système. L'objectif est de présenter ces méthodes, d'observer leurs résultats et de déterminer dans quelles situations elles sont applicables.

À propos du calcul de l'arbre des formes sur des images n -dimensionnelles en temps quasi-linéaire

par *Sébastien Crozet* 11 juillet, 10h00, Amphi 3

L'arbre des formes est une transformée d'image utile pour les traitements d'images discrètes de façon auto-duale. Dans un récent article nous avons présenté un nouvel algorithme de calcul de cet arbre en n dimensions. Cependant, aucune preuve ni étude de performances n'a été décrite. De plus, utilisé tel quel, l'algorithme nécessite une multiplication de la taille de l'image traitée telle que l'occupation mémoire et les temps de calcul sont très importants. Nous étudierons la preuve de l'algorithme ainsi que les détails d'initialisations. Nous apportons aussi une amélioration afin de réduire son occupation mémoire et ses temps de calcul lorsqu'il est appliqué à des images bidimensionnelles.

Speaker ID

Projection non-linéaire pour l'attribution de score selon la distance en cosinus dans le contexte des systèmes de vérification du locuteur à base d'i-vectors

par *Benjamin Roux* 11 juillet, 10h30, Amphi 3

À l'heure actuelle, l'espace des i-vectors est considéré comme le modèle standard pour la représentation d'informations et également grâce à l'utilisation de nouvelles méthodes de classification comme l'Analyse Discriminante Linéaire Probabiliste (PLDA) ou encore le classifieur à base de distance cosinus (CD). L'idée du scoring avec le CD est de projeter les caractéristiques évoluant dans un nombre de dimensions élevé sur une hypersphère de dimension plus faible. Aujourd'hui, tous les systèmes utilisent d'abord une