



The ImageCLEF 2013 Plant Identification Task

Hervé Goëau, Pierre Bonnet, Alexis Joly, Vera Bakić, Daniel Barthélémy,
Nozha Boujemaa, Jean-François Molino

► To cite this version:

Hervé Goëau, Pierre Bonnet, Alexis Joly, Vera Bakić, Daniel Barthélémy, et al.. The ImageCLEF 2013 Plant Identification Task. CLEF, Sep 2013, Valencia, Spain. hal-00960929

HAL Id: hal-00960929

<https://inria.hal.science/hal-00960929>

Submitted on 19 Mar 2014

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

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

The ImageCLEF 2013 Plant Identification Task

Hervé Goëau¹, Pierre Bonnet², Alexis Joly¹, Vera Bakic¹, Daniel Barthelemy³,
Nozha Boujemaa⁵, and Jean-François Molino⁴

¹ INRIA, IMEDIA & ZENITH teams, France, name.surname@inria.fr,
<http://www-rocq.inria.fr/imedia/>, <http://www-sop.inria.fr/teams/zenith/>

² CIRAD, UMR AMAP, France, pierre.bonnet@cirad.fr, <http://amap.cirad.fr>

³ CIRAD, BIOS Direction and INRA, UMR AMAP, F-34398, France,
daniel.barthelemy@cirad.fr, <http://amap.cirad.fr/fr/index.php>

⁴ IRD, UMR AMAP, France, jean-francois.molino@ird.fr, <http://amap.cirad.fr>

⁵ INRIA, Direction of Saclay Center, nozha.boujemaa@inria.fr,
<http://www-saclay.inria.fr>

Abstract. The ImageCLEFs plant identification task provides a testbed for a system-oriented evaluation of plant identification about 250 species trees and herbaceous plants based on detailed views of leaves, flowers, fruits, stems and bark or some entire views of the plants. Two types of image content are considered: *SheetAsBackground* which contains only leaves in a front of a generally white uniform background, and *NaturalBackground* which contains the 5 kinds of detailed views with unconstrained conditions, directly photographed on the plant. The main originality of this data is that it was specifically built through a citizen sciences initiative conducted by Tela Botanica, a French social network of amateur and expert botanists. This makes the task closer to the conditions of a real-world application. This overview presents more precisely the resources and assessments of task, summarizes the retrieval approaches employed by the participating groups, and provides an analysis of the main evaluation results. With a total of twelve groups from nine countries and with a total of thirty three runs submitted, involving distinct and original methods, this third year task confirms Image Retrieval community interest for biodiversity and botany, and highlights further challenging studies in plant identification.

Keywords: ImageCLEF, plant, leaves, leaf, flowers, fruits, bark, stem, species, retrieval, images, collection, identification, fine-grained classification, evaluation, benchmark

1 Introduction

Convergence of multidisciplinary research is a key to answer profound challenges of humanity related to health, biodiversity or sustainable energy. The integration of life sciences and computer sciences has a major role to play towards managing and analyzing cross-disciplinary scientific data at a global scale. More specifically, building accurate knowledge of the identity, geographic distribution

and uses of plants is essential if agricultural development is to be successful and biodiversity is to be conserved. Unfortunately, such basic information is often only partially available for professional stakeholders, scientists and citizens, and often incomplete for ecosystems that possess the highest plant diversity. A noticeable consequence, expressed as the *taxonomic gap*, is that identifying plant species is usually impossible for the general public, and often a difficult task for professionals, such as farmers or wood exploiters and even for the botanists themselves. The only way to overcome this problem is to speed up the collection and integration of raw observation data, while simultaneously providing to potential users an easy and efficient access to this botanical knowledge. In this context, content-based visual identification of plant's images is considered as one of the most promising solution to help bridging the taxonomic gap. Evaluating recent advances of the Image Retrieval community on this challenging task is therefore an important issue.

This paper presents the *plant identification task* that was organized for the third year running within ImageCLEF⁶ [10] dedicated to the system-oriented evaluation of visual based plant identification. Like previous year, the task is more related to a retrieval task instead of a pure classification task in order to consider a ranked list of retrieved species rather than a single brute determination. Visual content was being the main available information but with additional information including contextual meta-data (author, date, locality name and geotag, names at different taxonomic ranks) and some EXIF data. Each year try to take to the next level the challenge to a more realistic scenario by covering progressively one entire flora at the scale of one wide region like France. After two years focused exclusively on leaves mainly from Mediterranean tree species, the task focused this year on 250 species of herbs and trees species living in France with different views or *organs* of plants: photographs of flowers, fruits, barks, leaves and the entire view of the plants. Finally, it was two types of content which were considered: a *SheetAsBackground* category containing scans and scan-like photographs of leaves in a front of a generally white uniform white background, and a *NaturalBackground* with most of the time a cluttered natural background of the 5 types of organs. The main originality of this data is that it was specifically built through a citizen sciences initiative conducted by Tela Botanica⁷, a French social network of amateur and expert botanists. This makes the task closer to the conditions of a real-world application: (i) organs of the same species are coming from distinct plants living in distinct areas and with at distinct growing stages, (ii) pictures and scans are taken by different users that might not used the same protocol to collect the leaves and/or acquire the images, (iii) pictures and scans are taken at different periods in the year.

⁶ <http://www.imageclef.org/2013>

⁷ <http://www.tela-botanica.org/>

2 Task resources

2.1 The Pl@ntView dataset

Building effective computer vision and machine learning techniques is not the only side of the *taxonomic gap* problem. Speeding-up the collection of raw observation data is clearly another crucial one. The most promising approach in that way is to build real-world collaborative systems allowing any user to enrich the global visual botanical knowledge [15]. To build the evaluation data of ImageCLEF plant identification task, we therefore set up a *citizen science* project around the identification of common woody species covering the Metropolitan French territory. This was done in collaboration with Tela Botanica social network and with researchers specialized in computational botany.

Technically, images and associated tags were collected through a crowd-sourcing web applications [15], [13] and were all validated by expert botanists. Several cycles of such collaborative data collection and taxonomical validation occurred. Scans of leaves were the first type of pictures collected thanks to the work of active contributors from Tela Botanica since the summer 2009. The idea of collecting only scans of leaves first was to initialize training data with limited noisy background and to focus on plant variability rather than mixed plant and view conditions variability. This allowed to collect a first dataset of 2228 scans over 55 species. A first public crowd-sourcing web application⁸ was then opened in October 2010 and additional data were collected up to March 2011. The new collected images were either *scans*, or photographs with uniform background (referred as *scan-like* photos), or unconstrained *photographs* with natural background. It involved besides 15 new species from the previous set of 55 species. In April 2011 a new version of the web application has opened⁹ and the acquisition protocol was extended to 4 more types of views with a natural background mentioned below, and focusing to the same limited set of species. During the last two years, members from Tela Botanica contribute regularly every month on more and more species, the final ambition being to cover the entire vascular French flora (around 6000 species) with numerous pictures of different plant organs, with numerous plant observations spread all over France at different growing stages photographed by a crowd of photographers, introducing slowly over the months great visual and morphological variabilities. However, for each year task, we decided to limit the number of species in the task by keeping only the most populated ones in terms of images and plant observations. This is why like the first year we decided to focus again only on leaves during the ImageCLEF 2012 Plant Identification task with a number of 125 species, because at the time of the task we didn't collected sufficiently pictures of complementary organs. This year, we decided to propose to add these complementary views while we added to the previous dataset Pl@ntLeaves 125 new more species more focusing on herbaceous plants than trees in order to cover more diversity of the French

⁸ it is closed now, but a newer application can be found at <http://identify.plantnet-project.org/en/base/plantscan>

⁹ <http://identify.plantnet-project.org/fr/>

flora. Complementary views concerning flowers, fruits, stems and entire views associated to previous species and previous plant observations yet contained in the 2012 dataset were also added. Finally, the Pl@ntView dataset used within ImageCLEF2013 plant task contained 26077 images collected by 327 distinct contributors: 11031 for the *SheetAsBackground* category and 15046 for the *NaturalBackground* (in more details 16% of leaves, 18% of flowers, 8% of fruits, 8% of stems and 8% of entire plant). The figure 2.1 gives some examples illustrating the type of views, but illustrating also the fact that a species does not contain systematically at least one image for each organ.

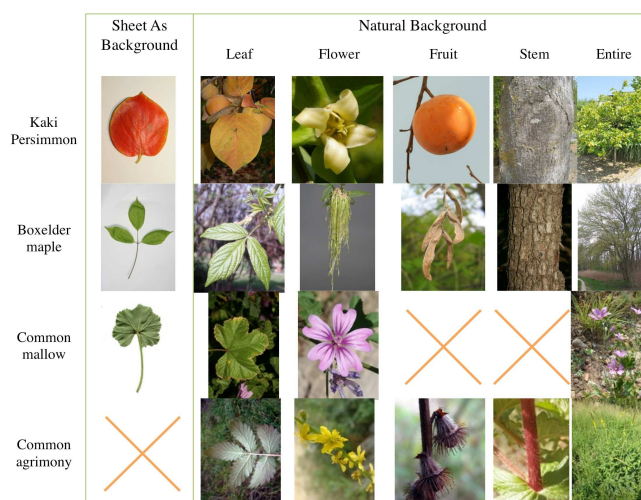


Fig. 1. Examples from the two categories and the five subcategories. Species does not contain systematically at least one image for each organ.

2.2 Pl@ntView metadata

Each image of Pl@ntView dataset is associated with the following meta-data:

- *IndividualPlantID*: plant observation identifier
- *Date*: date and time of plant observation
- *Type*: *SheetAsBackground* or *NaturalBackground*
- *Content*: *Flower*, *Fruit*, *Leaf*, *Stem* or *Entire*
- *Taxon*: full taxon name according the botanical database[4](Regnum, Class, Subclass, Superorder, Order, Family, Genus, Species)
- *ClassId*: species identifier
- *VernacularNames*: English common name
- *Author* name of the author of the picture
- *Organization* name of the organization of the author
- *Locality* locality name (a district or a country division or a region)
- *GPSLocality* GPS coordinates of the locality.

Concerning the locality information, note that sometimes the GPS can be very imprecise when the locality was not mentioned: in this case we used the GPS coordinates of the district or the country division or a region, according to the level of information available. Metadata is stored in independent xml files, one for each image. Additional but partial meta-data information can be found in the image's EXIF, and might include the camera or the scanner model, the image resolution and dimension, the optical parameters, the white balance, the light measures, etc.

2.3 About other plant datasets

A crucial added-value of this collection over older ones used in the literature (such as Swedish [22], ICL [1], Flavia [23] or Smithsonian [7]), is that it was built in a collaborative manner, through a citizen sciences initiative, and in collaboration with a well established social network specialized in botany. This makes it closer to the conditions of a real-world application: (i) pictures of organs of the same species are coming from distinct plants living in distinct areas (ii) pictures and scans are taken by different users that might not used the same protocol to collect the leaves and/or acquire the images (iii) pictures and scans are taken at different periods in the year. Intra-species visual variability and view conditions variability are therefore more stressed-out. In the end, this makes our identification challenge much more realistic but also more complex. We can mention here two other challenging datasets, the OxfordFlower[19] dataset, and the MobileFlora [5] one, which are indirectly built in a collaborative manner through web crawling but without (or very partially), contextual information like the author, the location, the date, etc. They unfortunately also come with a set of drawbacks or unrealistic properties: (i) they include only flower images (ii) they focus on the most represented species on the web rather than the most represented in a given area (iii) the definition of the taxonomic classes is not rigorous (sometimes genus, sometimes species, sometimes nothing well defined). Finally, the plant branch of the huge crowdsourced dataset ImageNet[12] could be interesting for our problem but it unfortunately contains too much errors, noisy classes and too sparse tags (typically about the type of view or the depicted organ).

2.4 Pl@ntViews variability

The ImageCLEF 2012 overview [14] provided numerous illustrations of the wide visual variability of the leaves. We present here more visual variability concerning the new introduced organs.

Flower There is a great diversity within flowers and it is an intensive subject of studies by botanists since the flower is often the key for identify a species. Flowers of the dataset can be categorized according to the color (see figure 2),

the symmetry (see figure 3), the number of petals (see figure 4) and the size (see figure 5). Most of the time one species is associated to one category, but there are exceptions like in figure 6 where for one same species the flowers can have different colors. Besides this first categorisation, botanists studied the



Fig. 2. Color variability of flowers.

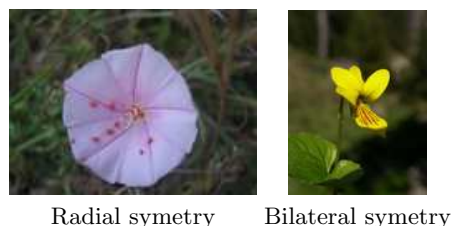


Fig. 3. Symetries of flowers.



Fig. 4. Number of petals.

inflorescence, i.e. the internal structure of the flower and the organisation of the flowers on a plant. Species from a same taxonomical group generally share a same organization, and thus a same visual appearance. Figure 7 gives all the type of inflorescence contained in the dataset. Some type of inflorescence can be very noticeable and very typical of a group of species. However, at the opposite some very distinct groups of species in the taxonomical hierarchy can have a very distinct visual appearance, but sharing a same type of inflorescence.

Fruit The fruit is the transformation of the flower and it can be also categorized into distinct types. The figure 8 shows the great diversity of type of fruits that



Fig. 5. Sizes of flowers.



Fig. 6. Color variability of flowers from one same species (*Iris lutescens* Lam.).



Fig. 7. Inflorescence types (structure of the flower(s) on the plant, how they are connected between them and within the plant).

are represented through the 250 species of the Pl@ntView dataset. The different modes of *dissemination* gives a second complementary and interesting way to show the visual diversity of the fruits in the dataset. Indeed, a same mode of dissemination of (even very) distinct species involves generally some same morphological features. For instance, for the *endozoochory* dissemination (seed dispersal by animals) the fruits are generally colored for attracting birds for instance.

Stem The stem is generally a difficult plant sub-part for identifying a species, maybe because the visual information is mainly expressed with the texture, less



Fig. 8. Fruit types.

Non-zoochory		Zoochory				
Anemonocory	Barochory	Dyszoochory	Endozoochory	Epizoochory	Mymezoochory	
(wind)	(gravity)	(rodents)	(excrements)	(fur, plumage)	(ants)	

Fig. 9. Examples of fruits according the dissemination. *Zoochory* involves animals in the seed dispersal.

by the color and or the shape. Age of the plant is second difficulty for analysing the stem, more precisely for the trees and theirs barks. Through the collaborative process, the dataset contains for numerous species, different plants at different ages and fill partially the wide diversity of the barks. The figure 10 shows a representative example for the species *Robinia pseudo-acacia* with young and old trees: more the tree is young more it has some thorns as a strategy of defence.



Fig. 10. The visual diversity of the bark of the *Robinia pseudoacacia*.

Entire The entire view is maybe the most difficult view for identifying with precision a species, because this kind of view generally does not contain sufficiently information, and because a same species can have a very different general appearance depending to the geographical and climatic conditions. However, more the plant is small (young or intrinsically small), more the "useful" organs for identification are visible. The figure 11 shows one big tree of *Magnolia grandiflora* L. where it is very difficult or even impossible for identifying the pant if we look the entire view. The second plant is a small herbaceous species of *Gentiana pneumonanthe* L. where we can see that the flower is very visible on the entire view for identification.

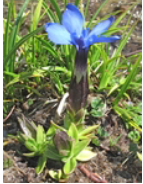
Species	Flower	Entire
<i>Magnolia grandiflora</i> L.		
<i>Gentiana pneumonanthe</i> L.		

Fig. 11. The entire views for herbaceous species and three are very different. The smaller one has the possibility to show useful organs for identification.

3 Task description

3.1 Training and Test data

The precise goal of the task was to retrieve the correct species among the top k species of a ranked list of returned species, one list for each image of a test dataset. Participants received a first training set of annotated images in order to explore different techniques and train their system. Six weeks later participants received the test set containing images without species labels, but with the view type, organ type, author, organization and plant identifier tags. Then, 2 months later, participants were allowed to submit up to 4 run files, most of the time related to variations of the same method. A particular attention was paid when splitting the data into training and test subsets to avoid any bias. Several pictures in the dataset might actually depict the same individual plant (or neighboring plants) observed in the same conditions (same person, day, device, lightening conditions, etc.). Randomly splitting images in a nave way would therefore favor having such near-duplicate images in both the training and the test subsets, making the recognition much more easy. To avoid this bias, we

therefore performed our random split at the observation level rather than at the image level thanks to associated metadata (observation id when available, author, date, etc.). Numerous images of the different views were automatically integrated in the training dataset since the associated plant observations were yet integrated last year task with the leaves. The training data finally resulted in 20985 images while the test data resulted in 5092 images. Detailed statistics of the composition of the training and test data are provided in Table 1.

		Images	Plants	Authors	Species
SheetAsBackground	Train	9781	732	36	126
	Test	1250	150	14	70
NaturalBackground	Train	11204	2553	176	244
	Test	3842	2454	229	238
Entire	Train	1455	955	104	234
	Test	694	567	107	177
Flower	Train	3521	1328	127	233
	Test	1233	970	142	203
Fruit	Train	1387	512	64	156
	Test	520	302	77	103
Leaf	Train	13285	1046	73	210
	Test	2040	420	68	143
Stem	Train	1337	629	38	131
	Test	605	408	35	77
All	Train	20985	11204	176	250
	Test	5092	3842	229	241

Table 1. Composition of the training and test data

3.2 Task objective and evaluation metric

According to similar concerns, the primary metric used to evaluate the submitted runs uses a two-stage average of raw image scores, one at the observation level (i.e. we compute the average score of all images belonging to the same observed plant), and one at the user level (i.e. we average the scores of the observations of a given user). A flat mean would actually have introduced some new bias with regard to a real world identification system. Indeed, as the dataset was built in a collaborative manner, it appears that few contributors often provide much more pictures than many other contributors who provided few (long tail distribution). Since we want to evaluate the ability of a system to provide correct answers to any user, we rather measure the mean of the average classification score per author. Furthermore, some authors sometimes provided many pictures of the same individual plant (to enrich training data with less efforts). Since we want to evaluate the ability of a system to provide the correct answer based on a single plant observation, we also decided to average the classification rate on each individual plant. The raw image score itself is computed for each test image as the inverse of the rank of the correct species in the list of retrieved species.

More formally, our primary metric was defined as the following average score S :

$$S = \frac{1}{U} \sum_{u=1}^U \frac{1}{P_u} \sum_{p=1}^{P_u} \frac{1}{N_{u,p}} \sum_{n=1}^{N_{u,p}} s_{u,p,n} \quad (1)$$

U : number of users (who have at least one image in the test data)

P_u : number of individual plants observed by the u -th user

$N_{u,p}$: number of pictures taken from the p -th plant observed by the u -th user

$s_{u,p,n}$: score between 1 and 0 equals to the inverse of the rank of the correct species for the n -th picture taken from the p -th plant observed by the u -th user

It is important to notice that while making the task more realistic, the normalized classification score also makes it more difficult. Indeed, it works as if a bias was introduced between the statistics of the training data and the one of the test data. It highlights the fact that bias-robust machine learning and computer vision methods should be preferred to train such real-world collaborative data.

Finally, to isolate and evaluate the impact of the image acquisition type (*SheetAsBackground*, *NaturalBackground*), a normalized classification score S was computed for each type separately. Participants were therefore allowed to train distinct classifiers, use different training subsets or use distinct methods for each data type.

4 Participants and techniques

With 12 finalist groups coming from all around the world over 9 countries and 33 submitted runs, the 2013 edition of the task confirmed its increasing attractiveness (respectively 10 and 11 groups crossed the finish line in 2011 and 2012) although its complexity was higher (with heterogeneous view types). Participants were mainly academics, specialized in computer vision and multimedia information retrieval. We list below the participants and give a brief overview of the techniques they used in their runs. We remind here that ImageCLEF benchmark is a system-oriented evaluation and not a deep or fine evaluation of the underlying algorithms. Readers interested by the scientific and technical details of any of these methods should refer to the ImageCLEF 2013 working notes of each participant (referenced below):

AGSPPR (3 runs) [25], China. AGSPPR team focused their work on the *SheetAsBackground* category and submitted 3 runs using distinct visual features and approaches: a global shape feature (i.e. the leafs area length of major axis and length of minor axis), SIFT features (run 2), and an extension of the CENTRIST approach (CENSus Transform hISTogram []) called SPACT designed for reducing the number of descriptors with a PCA algorithm (run 3). For run 1 and 3, they used a multiclass Support Vector Machine (SVM) classifier with a radial basis kernel function, while they used a pure matching approach for the 2nd run.

DBIS (4 runs) [20], Germany. DBIS team runs are based on global visual features and a multiclass SVM classifier. These participants experimented numerous (early) combinations of about thirty global features, in order to select the

best combination for each type of view. The selected features are predominantly based on color (Auto Color Correlogram, Border Interior Color, Color Histogram, Color Layout, Color Structure, EdgeHistogram, tamura, CEDD, FCTH). They also experimented several SVM parameters in order to boost their results.

I3S (2 runs) [16], France. I3S team used a popular approach in the field of image classification: they extracted SIFT features in order to produce Bag of visual Words (BoW), one BoW vector by picture, from a 1000 visual words dictionary (built with kMeans clustering algorithm). BoW's are then exploited to train species model with SVM classifiers, one for each species and type of view. Species prediction of test images are produced with a one-against-all procedure.

INRIA PLANTNET (4 runs) [8], France. For the *SheetAsBackground* category, after a basic Otsu segmentation, INRIA team used multiscale triangle representations, alone and combined with other shape-based descriptors (Directional Fragment Histogram and shape parameters). In addition, multi-image queries were considered, by using images belonging to the same plant observation in order to boost the results. For the *NaturalBackground* category, all the 4 submitted runs are based on local features (SURF, Fourier, rotation invariant Local Binary Patterns, Edge Orientation Histogram, weighted RGB and HSV histograms). The last one uses a multi-cue Fisher Vector embedding [1] with a one-against-all multiclass SVM classifier. The three first runs use Hamming embedding and hash-based approximate knn matching: all local features are hashed, indexed and searched in separate indices (one for each each type of view and type of feature) and retrieved images are scored by the number of matches. A two-stage late fusion scheme is then apply to combine the image response lists of the different modalities and of the different types of view. Metadata was also successfully used (in run 2), in particular the date for the flower category and the plant observation identifiers (to share the query images of the same plant).

LAPI (1 run) [9], Romania. LAPI team proposed to exploit a complex approach for image description based on contour extraction, curve partitioning and abstraction. They suggest that their approach is a "structural alternative" to the prevailing gradient-based features (e.g. SIFT). Contrary to other teams, they considered a more difficult task by automatically recognizing the view type before recognizing the plant species. They used a classical Linear Discriminant Analysis (LDA) as classifier for both the view type recognition and the species prediction.

LIRIS REVES (2 runs) [11], France. ReVes team used the same supervised model-based segmentation strategy than the one they used during the 2012 leaf-oriented campaign and tried to extend it to the other types of view (although it was more difficult to build a priori shape models of that organs). They used a late fusion approach to combine the decisions of the classifiers of each modality as well as to combine the multiple images of a given individual plant when this occurred in the query set. They finally attempted to use the geo-tags available in the metadata by interpolating them thanks to external climatic data.

MICA (3 runs) [17], Vietnam. MICA team experimented 3 distinct approaches. Run 1 used a GIST descriptor with a k-nearest neighbors rule on

all types of view. Run2 was based on the same approach but with additional color and texture features for the *Flower* and *Entire* types of view. Run3 used a Bag of visual Words (BoW) approach based on SURF local features and an "un-sharp masking" pre-processing step to filter some background information. Classification was achieved through a multi-class SVM.

NLAB UTOKYO (3 runs) [18], Japan. NLAB participant focused his work on visual features learning for building accurate image descriptions. A set of local features, mostly SIFT variations and a Self Similarity descriptor, were densely extracted in each picture according to a regular grid and then "augmented" with a supervised polynomial embedding technique taking into account neighboring local features. Further, these locally embedded and augmented features were encoded into a global Fisher Vector representation which allows an accurate classification with any linear classifier. In this work, linear logistic regression models were used. An independent classifier was trained for each raw descriptor and a late-fusion based an average log-likelihood of posterior probabilities was used to merge independent classifier results.

SABANCI-OKAN (1 run) [24], Turkey. This team submitted only one run using distinct features for the two categories. For the *SheetAsBackground* category, an automatic segmentation was performed using edge preserving morphological simplification by means of area attribute filters, followed by an adaptive threshold. Then, a variety of shape and texture features were extracted (the same than the ones used during the 2012 campaign). For the *NaturalBackground* category, a set of global features was extracted: HSV color auto-correlograms, weighed-saturation hue histogram and other texture descriptors, depending on the considered organ. For the *Flower*, *Fruit* and *Entire* view types only color features were used, while for the *Stem* view type, texture features were used after a segmentation preprocessing step. The dates provided in the metadata were also exploited for the three first view types (that are likely to be more time dependent). Classification was performed through independent SVM classifiers, one for each view type.

SCG USP (4 runs), Brazil. This team submitted one run with a fully automatic approach (run 1) and three other ones runs involving human assistance for a background/foreground segmentation. More precisely, training pictures were segmented with the semi-supervised Grabcut algorithm, while test images were manually segmented. Then, numerous features were extracted: Gabor, LBP, fractal, geometrical features. The final classification step was performed with a LDA classifier, except for the 3rd run where a SVM classifier was used. Only the 4th run tried to train independent classifiers (i.e. one for each view type).

UIAC (3 runs) [21], Romania. Unlike the other groups, UAIC explored the strategy of integrating additional external training data to boost their performances. They actually crawled 507 additional pictures from Wikimedia Commons with relevant annotations. And this confirms the difficulty of collecting dense and accurate data specific to a given flora. From the technological point of view, they used the LIRe (Lucene Image Retrieval) engine and, after preliminary tests, they selected the Joint Composite Descriptor (JCD). The LIRe

engines gives for each test image a list of training images where a candidate species potentially appears several times. Thus, they used 3 distinct approaches of combination in order to obtain a single score for each species: a max operator, a normalized sum, and a naive Bayes classifier. The results were further refined and ranked based on GPS metadata, author names and organization tags, assuming that certain authors and organizations would have a greater interest in certain plant species.

Table 2 attempts to summarize the methods used at different stages (feature, classification, subset selection,...) in order to highlight the main choices of participants. This table should be used in next section on result analysis, in order to see if there are some common techniques which tend to lead to good performances.

5 Results

5.1 Global analysis

We present here an overview of the official results of the task and discuss the main findings. **SheetAsBackground:** Table 3 and figure 12 present the identification scores of the 33 submitted runs for the *SheetAsBackground* category. As expected, results on scans and scan-like images of leaves are generally higher than the photographs of the *NaturalBackground* category. The Sabanci Okan teams reached the highest scores of 0.607 with an approach mainly centered on leaf shape boundary features. Using contour-based approaches is confirmed to be an effective strategy by the good performances of the Inria PlantNet group and the Liris team. Interestingly, one team which used a more generic approach in computer vision (the NLabUTokyo team working with Fisher Vector representations), also obtained very good identification scores whereas they used exactly the same technique for the *NaturalBackground* category. Other teams who attempted to use non-contour based approaches obtained significantly lower scores.

Compared to the raw identification scores obtained during the 2012 campaign, we only noticed a slight increase (0.607 vs 0.58 for scans and 0.55 for scan-like). But it is important to remark that the task itself was more complex in several aspects: (i) scans and scan-like pictures have been merged in a single category (ii) the number of species was increased from 115 (scans) or 83 (scan-like) to 126 this year (iii) test images themselves were more complex (weaker lighting conditions, more shadows, more old dried leaves and less uniform background).

NaturalBackground: Table 4 and figure 13 present the identification scores of the 33 submitted runs for the *NaturalBackground* category. As expected, results are significantly lower than the *SheetAsBackground* category due to the noisy backgrounds and clutter effects. The highest scores, obtained by the NLabU-Tokyo team, reached equivalent values than the 2012 task, but **without any human assistance** in the workflow, contrary to last year best runs that involved semi-automatic segmentation mechanisms. This is even more remarkable

images from a same Individual Plant during evaluation on training dataset.

Table 2. Approaches used by participants. In training subset column, *All* means that the 2 categories where not distinguished, while *Cat.* means it is the case, and *Subcat.* means that the sub-categories (*Flower*, *Frut*, *Leaf*, *Stem*, *Entire*) were considered for the *NaturalBackground* category. Column IP indicates if participants avoid to split

Team	Segmentation	Features	Classification	Train subsets	Metadata	IP
AGSPPR	✓ (run 1)	SPACT, SIFT, global shape	Multiclass SVM with a radial basis function kernel (run 1&3). Matching approach (run2).	<i>Cat.</i>		
DBIS	×	Auto Color Correlogram, Border Interior Color, Color Histogram, Color Layout, Color Structure, Edge Histogram, Tamura, CEDD, FCTH.	Distinct kernels by run for Multiclass SVM (RBF, polynomial, sigmoid)	<i>Subcat.</i>	×	×
I3S	×	SIFT + BOW	One against all Mutliclass SVM, linear kernel	<i>Subcat.</i>	×	×
IINRIA PLANTNET	Otsu (<i>Sheet</i> , and run 3 for <i>Natural</i> with an abort criterion	<i>Sheet</i> : Multiple triangular representations eventually combined with DFH descriptors and shape parameters <i>Natural</i> : Harris-like key points + SURF, LBP, weighted RGB, HSV, Fourier, EOH	Large Scale Matching approaches with late fusion schema or One against all multiclass SVM classifier applied to Fisher Vectors	<i>Subcat.</i>	Flowering date, Individual plant ids	×
LAPI	×	contour extraction with Canny + curve partitioning and abstraction + PCA	Linear Discriminant Analysis	<i>Subcat.</i>	×	✓
LIRIS REVES	Semi-supervised segmentation	Bagging features: Lab colors, Gabor wavelets, SURF Shape: centered moment, eccentricity, Hu/Zernike moments	Naive distance based classification	<i>Subcat.</i>	GPS with climate area extension, Individual plant ids	×
MICA	Auto	Gist (run1, 2), Color and texture histogram (run2 flower, entire), SURF + Bow (run3)	SVM	<i>Cat.</i> (?)	×	×
NLAB UTOKYO	×	dense grid SIFT, C-SIFT, Opponent-SIFT, HSV-SIF, self-similarity SSIM + polynomial embedding + Fisher Vector	Linear Logistic Regression, late fusion	<i>All.</i> , <i>Cat.</i> , <i>Subcat.</i>	×	✓
SABANCI-OKAN	Auto	<i>Sheet</i> : variety of contour based features, texture (Fourrier), color descriptors, edge background/foreground histogram. <i>Natural</i> : HSV color auto-correlograms, weigthed-saturation hue histogram,	SVM classifiers	<i>Subcat.</i>	date (flower, fruit)	✓
SCG USP	run1: auto, run 2-3-4: manual for test, semi for train	Gabor, LBP, fractal, geometrical	LDA or SVM (run3)	<i>Subcat.</i> (only run4)	GPS	×
UIAC	×	Joint Composite Descriptor	image result list fusion with max or sum operator, or naive bayesian approach	<i>All</i>	GPS, author, organization	×
VICOM-TECH	color segmentation	trace transform, shape relationship	Linear SVM 1-vs-all multi-class strategy	<i>All</i>	×	×

Run name	retrieval type	run-type	Score
Sabanci Okan Run 1	Visual	Auto	0,607
Inria PlantNet Run 2	Visual	Auto	0,577
Inria PlantNet Run 3	Visual	Auto	0,572
Inria PlantNet Run 1	Visual	Auto	0,557
Inria PlantNet Run 4	Visual	Auto	0,517
NlabUTokyo Run 1	Visual	Auto	0,509
NlabUTokyo Run 3	Visual	Auto	0,502
NlabUTokyo Run 2	Visual	Auto	0,502
Liris ReVeS Run 2	Mixed (textual + visual)	HA	0,416
Liris ReVeS Run 1	Mixed (textual + visual)	HA	0,412
Mica Run 3	Visual	Auto	0,314
DBIS Run 2	Visual	Auto	0,311
DBIS Run 4	Visual	Auto	0,281
LAPI Run 1	Visual	Auto	0,228
UAIC Run 4	Visual	Auto	0,205
DBIS Run 3	Visual	Auto	0,193
DBIS Run 1	Visual	Auto	0,191
AgSPPR Run 2	Visual	Auto	0,104
SCG USP Run 3	Textual	HA	0,103
UAIC Run 1	Visual	Auto	0,094
UAIC Run 2	Mixed (textual + visual)	Auto	0,088
UAIC Run 3	Visual	Auto	0,087
AgSPPR Run 1	Visual	Auto	0,071
AgSPPR Run 3	Visual	Auto	0,059
SCG USP Run 1	Textual	Auto	0,051
SCG USP Run 2	Textual	HA	0,051
I3S Run 1	Mixed (textual + visual)	Auto	0,039
I3S Run 2	Mixed (textual + visual)	Auto	0,039
SCG USP Run 4	Textual	HA	0,033
Mica Run 2	Visual	Auto	0,009
Mica Run 1	Visual	Auto	0,009
Vicomtech Run 1	Mixed (textual + visual)	Auto	0
Vicomtech Run 2	Mixed (textual + visual)	Auto	0

Table 3. Normalized classification scores for each run for the *SheetAsBackground*. HA=humanly assisted, Auto=full automatic.

given that their approach was purely based on the visual content contrary to the second best run of the task (by Inria Plantnet team) which did make use of the date and the plant identifier tags. The contribution of using the metadata can be observed by comparing this run with the second best one of that team (Inria Plantnet run 1) that was purely based on visual data. Overall, the runs of these two teams represent the head of the pack with six (or even seven) runs clearly outperforming the other runs.

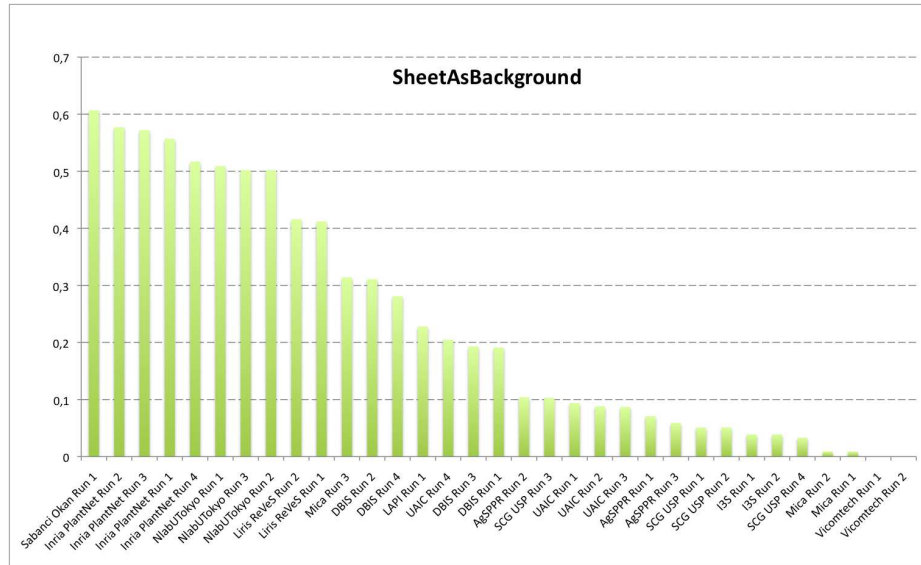


Fig. 12. Scores for *SheetAsBackground* category.

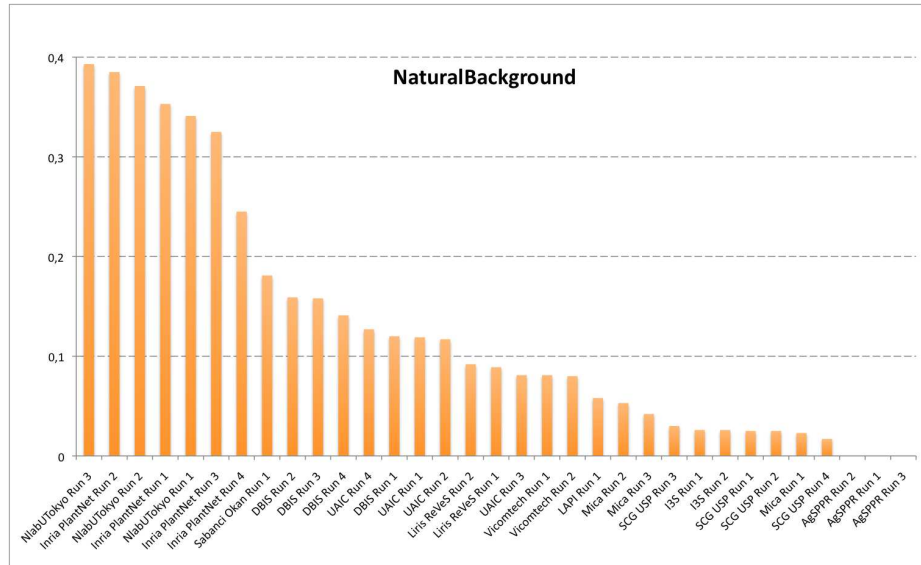


Fig. 13. Scores for *NaturalBackground* category.

Detailed results: Figures 17, 14, 18, 19, 16, and figure 15 display the identification scores for each view type separately (still for the *NaturalBackground* category). It shows that the average identification scores are significantly boosted by the good performances obtained on the flower images. Most techniques used

Run name	runfilename	Entire	Flower	Fruit	Leaf	Stem	Nat.
NlabUTokyo Run 3	run3	0,297	0,472	0,311	0,275	0,253	0,393
Inria PlantNet Run 2	plantnet inria run2	0,274	0,494	0,26	0,272	0,24	0,385
NlabUTokyo Run 2	run2	0,273	0,484	0,259	0,273	0,285	0,371
Inria PlantNet Run 1	plantnet inria run1	0,254	0,437	0,249	0,24	0,211	0,353
NlabUTokyo Run 1	all siftcophphsv cca	0,236	0,423	0,209	0,269	0,276	0,341
Inria PlantNet Run 3	plantnet inria run3	0,216	0,421	0,238	0,195	0,176	0,325
Inria PlantNet Run 4	plantnet inria run4	0,15	0,327	0,137	0,165	0,171	0,245
Sabancı Okan Run 1	Sabancı-Okan-Run1	0,174	0,223	0,194	0,049	0,106	0,181
DBIS Run 2	DBISForMaT run2 train2012 svm Scan12 Photo4 - 1 4	0,102	0,264	0,082	0,034	0,095	0,159
DBIS Run 3	DBISForMaT run3 cross- val2013 svm feature4 con- fig60 1 2 3	0,109	0,256	0,079	0,035	0,095	0,158
DBIS Run 4	DBISForMaT run4 cross- val2013 svm feature5 con- fig80 Photo14 1 3 3	0,152	0,206	0,104	0,027	0,042	0,141
UAIC Run 4	run wiki max 1	0,09	0,136	0,12	0,08	0,128	0,127
DBIS Run 1	DBISForMaT run1 train2012 svm Scan4 Photo2 1 2 3	0,067	0,168	0,1	0,052	0,103	0,12
UAIC Run 1	run wiki sum 3	0,089	0,109	0,132	0,093	0,104	0,119
UAIC Run 2	run author10 GSP10 lire80	0,092	0,105	0,127	0,096	0,11	0,117
Liris ReVeS Run 2	LirisReVeS run2	0,026	0,102	0,082	0,161	0,166	0,092
Liris ReVeS Run 1	LirisReVeS run1	0,021	0,098	0,081	0,151	0,153	0,089
UAIC Run 3	run lire naivebayes	0,068	0,055	0,111	0,049	0,102	0,081
Vicomtech Run 1	outputCLEFTestMean	0,095	0,117	0	0	0,1	0,081
Vicomtech Run 2	outputCLEFTestMax	0,091	0,116	0	0	0,094	0,08
LAPI Run 1	LAPI run1	0,026	0,073	0,025	0,084	0,043	0,058
Mica Run 2	MICA-run2	0,016	0,086	0,048	0,014	0,014	0,053
Mica Run 3	Run3	0,016	0,013	0,048	0,11	0,014	0,042
SCG USP Run 3	SCG USP run3	0,017	0,025	0,042	0,047	0,054	0,03
I3S Run 1	new 100	0,017	0,023	0,041	0,038	0,025	0,026
I3S Run 2	new2 100	0,017	0,023	0,041	0,038	0,025	0,026
SCG USP Run 1	SCG USP run1	0,02	0,026	0,027	0,02	0,037	0,025
SCG USP Run 2	SCG USP run2	0,027	0,029	0,02	0,018	0,019	0,025
Mica Run 1	MICA-run1	0,016	0,013	0,048	0,014	0,014	0,023
SCG USP Run 4	SCG USP run4	0,019	0,014	0,022	0,031	0,021	0,017
AgSPPR Run 1	AgSPPR run1	0	0	0	0	0	0
AgSPPR Run 2	AgSPPR run2	0	0	0	0	0	0
AgSPPR Run 3	AgSPPR run3	0	0	0	0	0	0

Table 4. Normalized and detailed cores for each run for the *NaturalBackground*.
HA=humanly assisted, Auto=full automatic.

by most participants were significantly more accurate on that image type. This confirms the botanical expertise on the important role of flowers in the identification mechanisms as this organ was historically used as the primary one to distinguish species between each others (for flowering plants of course). This is good news that computer vision methods go in the same direction.

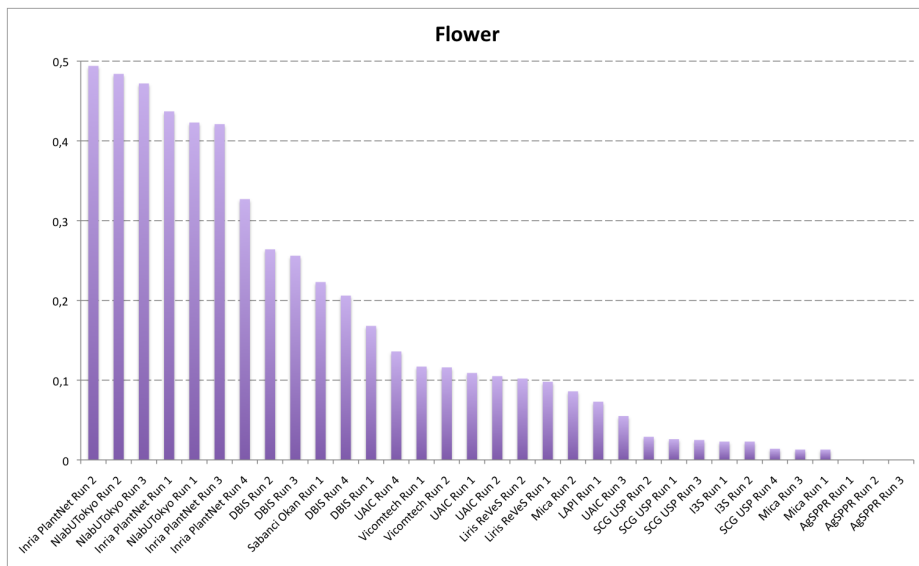


Fig. 14. Detailed scores for *Flower* subcategory.

Besides the *Flower* category, there was no clear second best organ or view type. *Stem* images provided surprisingly good results relatively to the botanist knowhow. Bark morphology is actually not considered as a the most accessible identification criterion for non-specialists. The texture itself is for instance highly correlated with the age of the plant. Identification results on the *Entire* plant views are also rather surprising regarding their higher complexity and variability. Overall, an important remark is that the ranking of the runs did not change much from an organ to another one, fostering the idea that generic methods might solve heterogeneous fine-grained classification problems.

Metadata: Regarding the use of metadata, two runs (Sabancı Okan run 1 and Inria Plantnet run 2) exploited successfully the date for improving the results. Using the observation date complementary to the visual content was a simple and efficient way to obtain a gain of up to 4 points on the *Flower* view type (thanks to the relatively short flourishing period of many species). Inria Plantnet run number 2 exploited also the observation identifier tag in order combine the results of the query images coming from the same plant. But since the whole *NaturalBackground* test dataset did contain only a few plant observations with multiple images, the impact of using this tag is much lower than the impact of

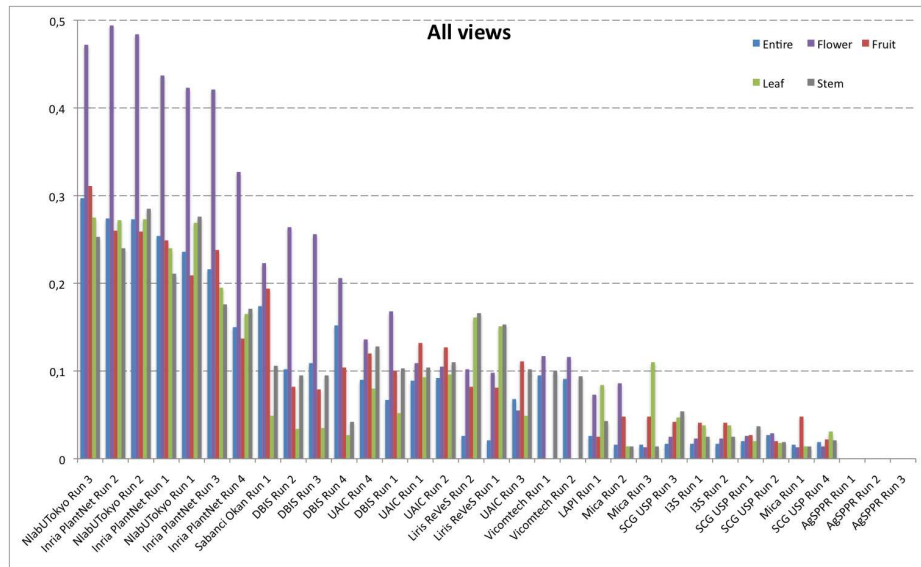


Fig. 15. Detailed scores for *NaturalBackground* subcategories.

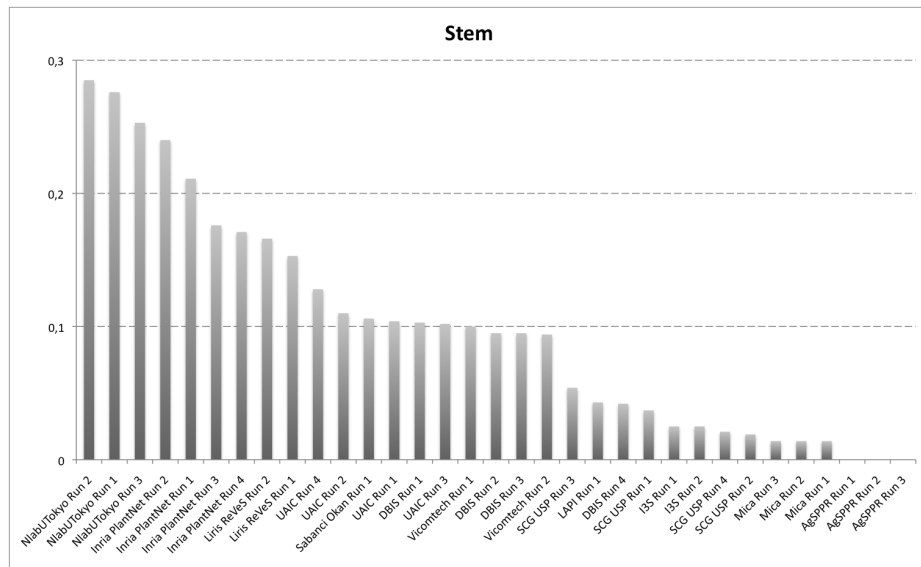


Fig. 16. Detailed scores for *Stem* subcategory.

using the date field. On the other side, this multiple-image strategy was much more beneficial for the *SheetAsBackground* category as a significant number of plants were represented by several images (leaves used for scans are actually more likely to be collected in mass from the same plant). The runs of Inria

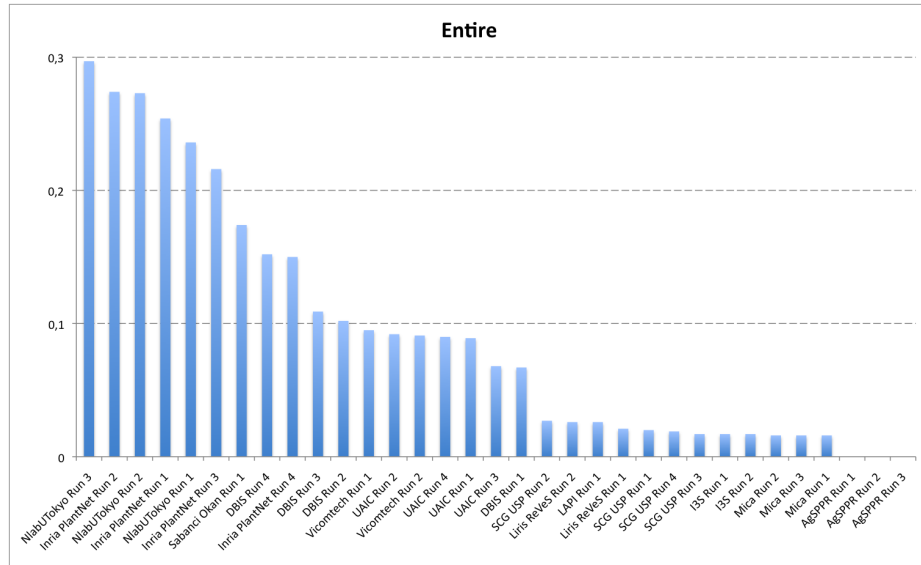


Fig. 17. Detailed scores for *Entire* subcategory.

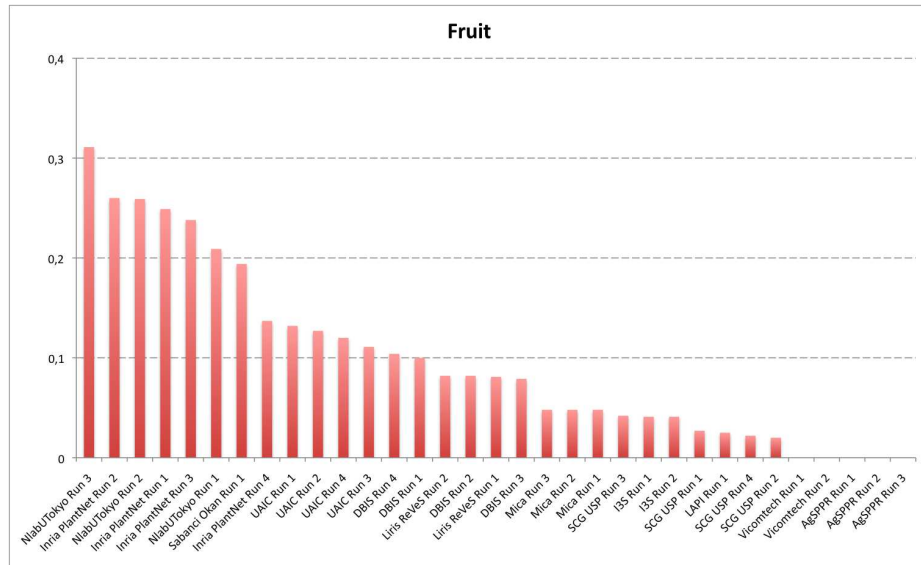


Fig. 18. Detailed scores for *Fruit* subcategory.

Plantnet and Liris ReVes teams exploited successfully this information for the *SheetAsBackground* category.

As the previous years, several teams, like Liris Reves or UIAC, attempted to exploit the geo-localization information in order to refine candidate species list.

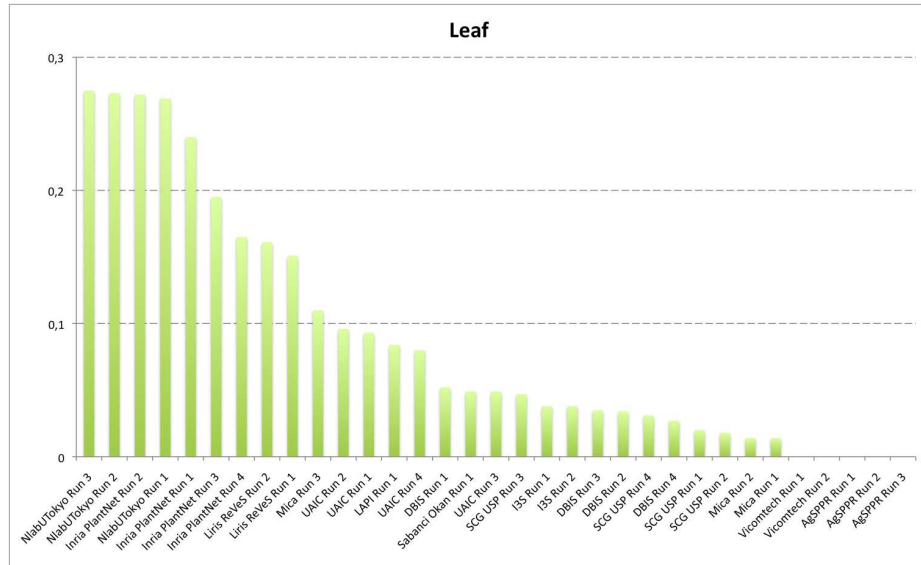


Fig. 19. Detailed scores for *Leaf* subcategory.

In particular, Liris teams proposed to use the raw GPS data of the training set complementary to external environmental data in order to interpolate them and build coarse species distribution maps. These maps were used afterwards to prune the species returned by the visual search and keep only the most probable ones. Unfortunately, the results do not show a great improvement over the purely visual runs of these teams. This can be explained by the fact that the database doesn't yet contain enough numerous and dense observations to build an accurate geographic repartition of the species. Also, the geo-localization data is partially noisy due to heterogeneous precisions in the localization (points, cities, departments).

Finally, the UAIC team tried to explore author and organization tags assuming that an authors or a group of authors from a same organizations have more interest on specific groups of species. However the results did not show clearly some gain by using these user informations. None of teams neither explored the hierarchical taxonomy structure, nor the common names, which could be source of improvements.

External data: UAIC explored the strategy of integrating additional external training data to boost their performances. They focused their search on Wikimedia Commons which contain more and more reliable contents related to species of life in general. They managed to crawl 507 additional pictures which is fine but clearly not sufficient to make a strong difference compared to tens of thousands of images in the training set. This confirms the difficulty of collecting dense and accurate data, specific to a given flora, and with relevant annotations (like organ and view type).

Impact of the global training strategy: Whereas some of the teams used a classical leave-one-image-out strategy cross-validate their training, some other ones used a more sophisticated leave-one-plant-out strategy that is closer to the real-world problem evaluated by the task. This second option seems to have taken benefits to the teams using it, namely Sabanci Okan, NLabUTokyo, Liris ReVes and Inria Plantnet. Indeed, they all mentioned that they did not split images from the same individual plant in the training set, in order to avoid overfitting problems (images of the same plant can actually be very similar).

Back to purely visual approaches: I3S and MICA teams experimented, at least through one run, a standard approach in image categorization with SIFT or SURF features, visual bag of words (BoW) representations and SVM multiclass. MICA team obtained intermediate scores contrary to I3S team. Explanations for this difference, can be that MICA use of preprocessing step for unsharp mask of *Leaf* images from *SheetAsBackground* and *NaturalBackground* (see figure12 MICA run 3 and I3S runs where scores are very different on *Leaf*). The more recent approach in image categorization based on Fisher Vector (FV) representations, which can be seen as an extension of BoW, showed a clear gain regarding to the BoW runs as we can see with the 3 NLabUTokyo runs and the Plantnet Inria run 4 on the *NaturalBackground* category. NLabUTokyo obtained the best scores, maybe because they capture local spatial information by enriching dense local descriptors with polynomials, contrary to the Inria Plantnet run where patches are extracted around Harris corners and descriptors are directly embedded in Fisher Vector representations. Moreover NLabUTokyo used also a late fusion where classifiers are trained independently for each descriptor, while Inria Plantnet run 4 used an intermediate fusion by concatenating FV representations from the different type of descriptors. Besides, late fusion is also a shared approach for the best runs of Inria Plantnet team.

The fact that NLabUTokyo runs obtained almost the best results for all subcategories, confirms the idea that FV representation is a successful generic approach in spite of different type of visual contents. It is important to notice that the run 2 obtained close scores to the best one (run 3) without considering subcategory tags, which show that views tags are maybe not essential for succeeding the task. This is an important conclusion since image tagging is an heavy process with users. However, this generic approach is not the most efficient on *SheetAsBackground* compared to contour based approaches dedicated to leaf shape analysis. This may show that generic approaches like the one used by the NLabUTokyo team is dependent to the background, and that through a dense grid patch extraction, their system learn a contextual information off the background. In particular this can be observed with the *Fruit* subcategory where NLabUTokyo run 3 outperforms other methods: fruits are generally small elements in the pictures difficult to capture, and also these organs appear often after the leafage, thus we can suppose these cluttered backgrounds have a non negligible contribution in the species contribution.

5.2 Performances per morphological features

Like in the previous working note with the leaf [14], we try here to present some complementary results by analysing some morphological features, more precisely on the sexual organs which are the flower and the fruit. Beyond the methods used, we try to analyse which feature, which kind of flower or fruit is intrinsically more difficult than the others. The figure 20 shows detailed results by category of color. The graph to the left shows the proportion of image test by color used for computing and displaying the graph to the right. Results in this second graph are sorted in a decreasing order of mean performance over all the submitted run (except the AgSPPR's runs which not really participate to the *NaturalBackground*). One can notice that the two most represented colors, the yellow and the white (more than 50% of the database) are not the ones which enables the best results, maybe rightly because there is more species and thus more confusions and ambiguity. Green flowers, which is not so rare, seem to be the most difficult color maybe because it expresses no color in a sense and thus it is a difficult information to capture, notably if the flowers hidden with a background of leafage or grass. Similarly, the brown flowers may also be very confused with barks for trees where flowers appear before leaves, which can explain the performances on this color. The figure 21 attempts to give a

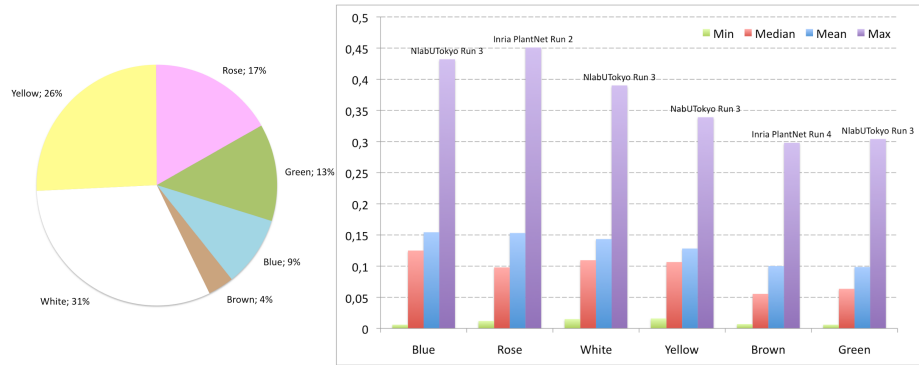


Fig. 20. Detailed results by flower color. The graph to the left represent the proportions of the tested images used for computing the detailed results in the graph to the right. This second graph gives the minimum, the median, the mean and the maximum scores over all the submitted runs for each flower color.

complementary perspective of results about flower according to the *inflorescence* structure as mentioned in section 2.4. First of all we can see that the inflorescence categories are strongly unbalanced in terms of image number. Thus the best scores obtained by *Cyathium*, *Panicle*, *Umbel* and to a lesser extent *Umbel*, *Capitulium* and *Solitary* are no very representative for making some relevant conclusions on these types of inflorescences. Concerning the most representative ones the *Cyme* seems to be the type where the runs performed the best on

average. The figure 22 gives the results according to the fruit type. As for the

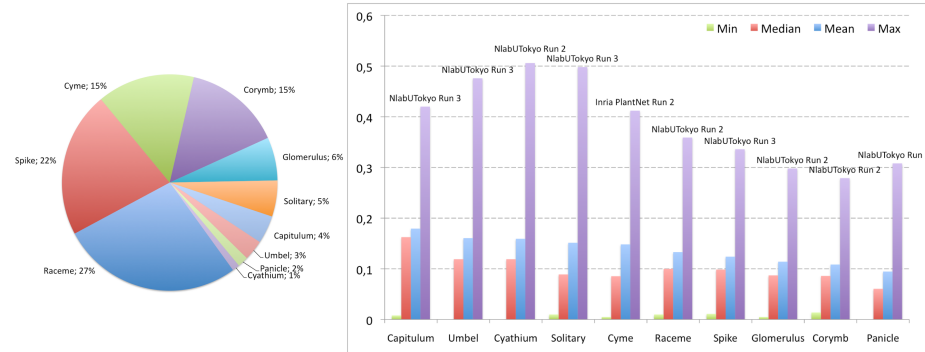


Fig. 21. Detailed results by inflorescence type

inflorescence, unfortunately, some types of fruits are not well represented in the dataset like *Silique*, *Cone*, *Folicle* and to a lesser extent *Legume*. Even if this last type of fruit is not so well represented, it is interesting to note that all teams seem to have the best scores on *Legume* because the associated species are from a very large family of plants called *Fabaceae* which is spread all over the world. Then, it is difficult to highlight one type of fruit over the others, because there is always one best method at the same score around 0.4. We have just to note that the *Samara* (like "helicopters" from maple for instance) seem to be clearly the most difficult type of fruit, even when we look at the best run (not over 0.2).

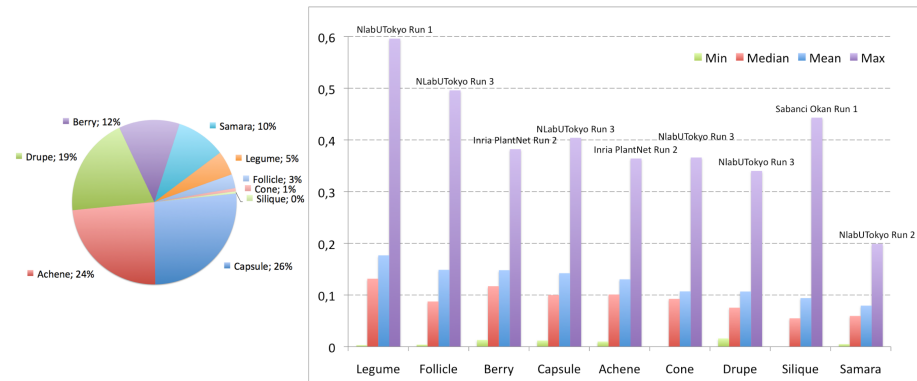


Fig. 22. Detailed results by fruit type.

6 Conclusions

This paper presented the overview and the results of ImageCLEF 2013 plant identification testbed following the two previous one in 2011 and 2012. The number of participants increased from 8 to 12 groups showing an increasing interest in applying multimedia search technologies to environmental challenges. This year the challenge climb one step by considering multiple type of view and organs of plants while the number of species increased from 125 to 250 species and plant observations densely covered the French territory. Results are encouraging by scaling state-of-the-art plant recognition technologies to a real-world application with thousands of species might still be a difficult task. Despite increasing difficulties on *SheetAsBackground* images and the number of species, scores are high and show that leaf analyses is still the best way for identifying a plant, even if collecting new scans is more difficult than shooting photographs. Performances obtained on *NaturalBackground* category of unconstrained pictures of plant organs are very encouraging especially for the *Flower* when we look detailed results and where best methods can compete with scores *SheetAsBackground*. It corroborates a well-know usage of botanists for identifying plants and this is good news in a sense that computer vision methods go in the same direction. An interesting conclusion is that these good results on *NaturalBackground* images are obtained with generic visual classification technique without any specificity related to plants. With the emergence of more and more plant identification apps [2] [6], [5], [3] and the ecological urgency to build real-world and effective identification tools, we believe that the detailed results and conclusions of the task will be of high interest for the computer vision and machine learning community.

Acknowledgements

This work was funded by the Agropolis fundation through the project Pl@ntNet (<http://www.plantnet-project.org/>) and the EU through the CHORUS+ Coordination action (<http://avmediasearch.eu/>). Thanks to all participants. Thanks to Jennifer Carré, Violette Roche and all contributors from Tela Botanica. Thanks to Souheil Selmi from Inria and Julien Barbe from Amap for their help.

References

1. The icl plant leaf image dataset, <http://www.intelengine.cn/English/dataset>
2. Leafsnap (May 2011), [https://itunes.apple.com /fr/app/leafsnap/id430649829](https://itunes.apple.com/fr/app/leafsnap/id430649829)
3. Folia (Nov 2012), " [https://itunes.apple.com /fr/app/fofia/id547650203](https://itunes.apple.com/fr/app/fofia/id547650203)
4. Tropicos (August 2012), <http://www.tropicos.org>
5. Mobile flora (May 2013), [https://itunes.apple.com /us/app/mobileflora/id592906385](https://itunes.apple.com/us/app/mobileflora/id592906385)
6. Plantnet (2013), [https://itunes.apple.com /fr/app/plantnet/id600547573](https://itunes.apple.com/fr/app/plantnet/id600547573)

7. Agarwal, G., Belhumeur, P., Feiner, S., Jacobs, D., Kress, J.W., R. Ramamoorthi, N.B., Dixit, N., Ling, H., Mahajan, D., Russell, R., Shirdhonkar, S., Sunkavalli, K., White, S.: First steps toward an electronic field guide for plants. *Taxon* 55, 597–610 (2006)
8. Bakić, V., Mouine, S., Ouertani-Litayem, S., Verroust-Blondet, A., Yahiaoui, I., Goëau, H., Joly, A.: Inria’s participation at imageclef 2013 plant identification task. In: Working notes of CLEF 2013 conference (2013)
9. C., R., L., F., C., V.: Has an image classification approach any chance at all (in plant classification)?... In: Working notes of CLEF 2013 conference (2013)
10. Caputo, B., Muller, H., Thomee, B., Villegas, M., Paredes, R., Zellhofer, D., Goëau, H., Joly, A., Bonnet, P., Gomez, J.M., Varea, I.G., Cazorla, M.: ImageCLEF 2013: the vision, the data and the open challenges. In: Proc CLEF 2013. LNCS (2013)
11. Cerutti, G., Tougne, L., Sacca, C., Joliveau, T., Mazagol, P.O., Coquin, D., Vacavant, A.: Late information fusion for multi-modality plant species identification. In: Working notes of CLEF 2013 conference (2013)
12. Deng, J., Dong, W., Socher, R., Li, L.J., Li, K., Fei-Fei, L.: ImageNet: A Large-Scale Hierarchical Image Database. In: CVPR09
13. Goëau, H., Bonnet, P., Barbe, J., Bakic, V., Joly, A., Molino, J.F., Barthélemy, D., Boujemaa, N.: Multi-organ plant identification. In: Proceedings of the 1st ACM international workshop on Multimedia analysis for ecological data. MAED ’12 (2012)
14. Goëau, H., Bonnet, P., Joly, A., Yahiaoui, Itheri Boujemaa, N., Barthélémy, D., Molino, J.F.: The ImageCLEF 2012 plant identification task. In: ImageCLEF (2012)
15. Goëau, H., Joly, A., Selmi, S., Bonnet, P., Mouysset, E., Joyeux, L.: Visual-based plant species identification from crowdsourced data. In: Proceedings of ACM Multimedia 2011 (2011)
16. Issolah, M., Lingrand, D., Precioso, F.: Sift, bow architecture and one-against-all support vector machines. In: Working notes of CLEF 2013 conference (2013)
17. Le, T.L., Pham, N.H.: Imageclef2013 plant identification mica. In: Working notes of CLEF 2013 conference (2013)
18. Nakayama, H.: Nlab-utokyo at imageclef 2013 plant identification task. In: Working notes of CLEF 2013 conference (2013)
19. Nilsback, M.E., Zisserman, A.: A visual vocabulary for flower classification. In: CVPR06
20. Saretz, S., Böttcher, T.: Btu dbis’ at imageclef2013 plant identification task. In: Working notes of CLEF 2013 conference (2013)
21. Serba, C., Siriteanu, A., Gheorghiu, C., Iftene, A., Alboaie, L., Breaban, M.: Combining image retrieval, metadata processing and naive bayes classification at plant identification 2013. In: Working notes of CLEF 2013 conference (2013)
22. Söderkvist, O.J.O.: Computer Vision Classification of Leaves from Swedish Trees. Master’s thesis, Linköping University, SE-581 83 Linköping, Sweden (September 2001), liTH-ISY-EX-3132
23. Wu, S.G., Bao, F.S., Xu, E.Y., xuan Wang, Y., fan Chang, Y., liang Xiang, Q.: A leaf recognition algorithm for plant classification using probabilistic neural network (2007)
24. Yanikoglu, B., Aptoula, E., Yildiran, S.T.: Sabanci-okan system at imageclef 2013 plant identification competition. In: Working notes of CLEF 2013 conference (2013)
25. Zhang, L., Cai, C.: Agsppr at imageclef 2013 plant identification task. In: Working notes of CLEF 2013 conference (2013)