

- Kassian A., Starostin G., Dybo A. & Chernov V., 2010, The Swadesh wordlist: an attempt at semantic specification, *Journal of Language Relationship* 4, p. 46-89.
- Klimov G. A., 1964, *Etimologičeskij slovar' kartvel'skich jazykov [Etymological dictionary of Kartvelian languages]*, Moscow, Izdatel'stvo Akademii Nauk SSSR.
- Kogan L., 2006, Lexical Evidence and the Genealogical Position of Ugaritic (I), *Babel und Bibel* 3, *Orientalia et Classica* vol. XIV, Winona Lake, Indiana, Eisenbrauns, p. 429-488.
- Kurilov G. N., 2001, *Jukagirsko-russkij slovar' [Yukaghir-Russian dictionary]*, Novosibirsk, Nauka Publishers.
- List J.-M. & Nelson-Sathi S. & Geisler H. & Martin W., 2014, Networks of lexical borrowing and lateral gene transfer in language and genome evolution, *BioEssays* 36.2, p. 141-150.
- Nikolayev S. & Starostin S., 1994, *A North Caucasian Etymological Dictionary*, Moscow, Asterisk Publishers.
- Nikolayeva I., 2006, *A Historical Dictionary of Yukaghir*, Berlin – New York, Mouton de Gruyter.
- Orel V., 2003, *A Handbook of Germanic Etymology*, Leiden – Boston, Brill.
- Pokorny J., 1958, *Indogermanisches etymologisches Wörterbuch*, Bern-München, Francke Verlag.
- Rastorguyeva V. S. & Edelman D. I., 2007, *Etimologičeskij slovar' iranskikh jazykov. Tom 3: f-h [Etymological dictionary of Iranian languages. Vol. 3: f-h]*, Moscow, Vostochnaja literatura.
- Starostin G., 2010, Preliminary lexicostatistics as a basis for language classification: A new approach, *Journal of Language Relationship* 3, p. 79-116.
- Starostin G., 2013a, Lexicostatistics as a basis for language classification: increasing the pros, reducing the cons, in H. Fangerau, H. Geisler, Th. Halling, W. Martin (ed.), *Classification and Evolution in Biology, Linguistics and the History of Science: Concepts – Methods – Visualization*, Stuttgart, Franz Steiner Verlag, p. 125-146.
- Starostin G., 2013b, *Jazyki Afriki. Opyt postrojenija leksikostatisticheskoy klassifikacii. Tom I: Metodologija. Kojanskije jazyki [Languages of Africa: an attempt at a lexicostatistical classification. Volume I: Methodology. Khoisan Languages]*, Moscow, Jazyki slavjanskoy kultury.
- Starostin S., 1991, *Altajskaja problema i proiskozhdenije japonskogo jazyka [The Altaic problem and the origins of the Japanese language]*, Moscow, Nauka Publishers.
- Starostin S., 1995, *Sravnitel'nyj slovar' jenijskich jazykov [Comparative dictionary of Yeniseian languages]. Ketskij sbornik. Lingvistika [Ket Volume. Linguistics]*, Moscow, Vostochnaja literatura, p. 176-315.
- Turner R. L., 1966, *A Comparative Dictionary of the Indo-Aryan Languages*, Oxford University Press.
- Vries J. de, 1962, *Altindisches Etymologisches Wörterbuch*, Leiden, Brill.

ReFLex : la reconstruction sans peine

Guillaume Segerer*

INTRODUCTION

Cet article se propose de décrire le fonctionnement d'un module d'aide à la reconstruction phonologique et lexicale. Ce module est intégré au site web d'exposition et d'exploitation de ReFLex (www.reflex.cnrs.fr), projet dont l'auteur de ces lignes est à l'initiative et qui s'articule autour d'une base de données lexicales sur les langues d'Afrique. Par conséquent, les exemples concrets cités ci-dessous sont empruntés aux langues africaines, mais l'ensemble d'outils dont il va être question peut être utilisé pour n'importe quel ensemble de langues.

Je présenterai d'abord le projet ReFLex dans son ensemble, puis je décrirai avec davantage de détails le module d'aide à la reconstruction (que l'on a tendance à appeler plus brièvement 'module de reconstruction' bien que la machine ne prenne pratiquement aucune décision). Ce module est, et sera encore probablement longtemps, en cours de perfectionnement. Certaines des caractéristiques décrites ici sont appelées à évoluer rapidement pour offrir encore plus de confort d'utilisation, et j'utilisateuse qui voudrait jouer à reconstruire après la lecture de ce texte trouvera très certainement des fonctionnalités nouvelles. Je tenterai ici d'évoquer les nouveautés en cours de développement.

1. REFLEx : PRÉSENTATION GÉNÉRALE

1.1. Les principes

ReFLex est un projet combinant une base de données lexicales sur les langues d'Afrique, des outils de traitement et un site internet pour l'accès aux données et aux outils¹. Le projet a réellement démarré en 2010, grâce à un financement de

* ILACAN-Sorbonne Paris Cité-INALCO – CNRS. Courriel : guillaume.segerer@cnrs.fr. Ce travail est lié au programme "Investissements d'Avenir" géré par l'Agence Nationale de la Recherche ANR-10-LABX-0083 (Labex EFL).

Avertissement : Cet article décrit une application en ligne, et par conséquent contient inévitablement des images ou 'captures d'écran' dont les couleurs peuvent être difficiles à discerner dans une version en niveaux de gris. Ces images ne peuvent donc fournir qu'une illustration approximative du contenu de l'écran. Le lecteur est donc encouragé à consulter le site original.

¹ Le développement a été assuré par Sébastien Flavie, ingénieur d'études au laboratoire DDL (Lyon), partenaire du projet

l'ANR. Le financement a pris fin en 2015 mais le projet reste un des axes forts du programme de recherche du laboratoire LLACAN.

Ce projet est unique car il associe plusieurs principes qui font généralement défaut dans les bases de données lexicales en ligne :

- il s'agit véritablement d'un lexique de référence, dans la mesure où les données consultables sont liées à des images numériques des sources originales dont elles proviennent. C'est cette caractéristique unique qui a donné son nom au projet (ReflEx : Référence Lexicon). Ainsi par exemple, pour la première fois, les mesures lexicostatistiques (un des exercices favoris de la linguistique comparative des langues africaines) deviennent intégralement reproductibles, et les choix concernant les ressemblances lexicales peuvent être discutés.

- Beaucoup des manipulations que l'on peut souhaiter faire sur les données sont réalisables en ligne : recherches sur critères complexes, comptages divers, statistiques combinatoires (voir K. Pozdniakov, ce volume), représentation cartographique des données, et surtout assistance à la reconstruction, module qui fait l'objet de cet article.

1.2. Les données : langues

Pour associer des données lexicales à une langue, il est nécessaire de disposer d'un inventaire des langues d'Afrique. Plusieurs inventaires existent actuellement, les deux principaux étant le catalogue correspondant à la norme iso 639-3, disponible sur le site www.ethnologue.com, et celui de l'institut Max Planck de Leipzig, récemment publié (<http://glottolog.org/>). Ces deux catalogues visent non seulement à inventorier toutes les langues du monde, mais en proposent en outre une classification généalogique. Les deux classifications sont assez différentes, celle de Glottolog étant beaucoup plus prudente (autrement dit, elle comporte davantage de familles). Notre catalogue (disponible à <http://reflex.cims.fr/Lexiques/webcat/index.html>) est différent des deux catalogues cités ci-dessus. Du fait de la présence au LLACAN de spécialistes reconnus de plusieurs familles de langues africaines, il est plus détaillé pour certaines zones géographiques, notamment l'Afrique centrale, la région Sénégal-Guinée, certaines parties du Nigeria. En septembre 2015, il recense 2441 langues.

1.3. Les données : sources

Comme ReflEx fournit des données lexicales provenant de documents publiés (pour la grande majorité, mais voir ci-dessous), avec accès aux documents d'origine, chaque lexique, en plus d'être associé à une langue, est associé à un document. Cette combinaison d'une référence linguistique et d'une référence bibliographique est appelée source dans ReflEx, ce qui correspond en partie au concept de *doculect*² introduit par Jeff Good en 2006. Par exemple, le document dont l'identifiant est 18737 est un document de Keith Snider publié au Ghana en 1989 et contenant des listes lexicales pour 5 langues du groupe Kwa, auxquelles

correspondent donc autant de sources dans ReflEx. Inversement, la langue joola karon (Atlantique) est représentée dans ReflEx par 4 documents : Barry 1987, Carlton & Rand 1994, Sambou 2007 et Wilson 2007, ce qui fait donc 4 sources.

Une source de ReflEx n'est pas la combinaison unique d'une langue et d'une référence bibliographique. En effet, lorsqu'un document contient des informations à propos de variantes dialectales, il n'est pas toujours facile de décider si ces variantes doivent avoir le statut de langue. Par exemple, Fresco 1970 (réf. 6910) fournit des listes lexicales pour le yoruba commun ainsi que pour 7 variétés dialectales de yoruba. Ces 8 sources sont toutes la combinaison de la référence n° 6910 (Fresco 1970) et de la langue n° 835 (yoruba).

Actuellement (juin 2016), le nombre de sources dans ReflEx est de 1164 et correspond à 705 langues.

1.4. Les données : lexiques

Toutes les données lexicales des langues d'Afrique ont par nature vocation à être intégrées à ReflEx. Pour des raisons pratiques cependant, certaines critères peuvent influencer sur le choix des données à importer en priorité :

- Les données publiées sont préférées aux données brutes, pour des questions de stabilité. Cependant, dans certains cas, notamment lorsqu'il s'agit de données de première main concernant des langues peu ou pas documentées, des données non publiées peuvent être incluses dans ReflEx.

- Les données déjà numérisées (en partie ou en totalité), ou celles déjà saisies sont évidemment plus faciles à intégrer.

- Les données qui font l'objet d'un travail de recherche particulier sont également prioritaires. Le projet ANR ReflEx comportait des tâches précises et impliquait une vingtaine de chercheurs. Nous avons donc mis l'accent sur les données concernées par ces tâches.

La taille des lexiques est très variable. Les plus petits ont quelques dizaines d'entrées, les 3 plus importants ont plus de 20000 entrées chacun (il s'agit du bambara, du peul du Maasina et du sereer).

- Leur structure est également très variable et va de la simple liste de mots au dictionnaire complexe. Certaines informations, comme les exemples d'illustrations, ne sont pas reprises dans ReflEx. En revanche, les formes citées dans les sources servent de base pour extraire de l'information nouvelle : schéma tonal, schéma syllabique, code phonétique. Ces informations vont permettre d'effectuer des recherches et des tris qui peuvent se révéler très utiles lors de la recherche de cognats.

Par ailleurs, les deux champs indispensables que sont la forme et le sens font l'objet de procédures d'harmonisation. Pour la forme, il s'agit surtout d'unifier les transcriptions de façon à pouvoir facilement repérer des formes semblables. En effet, selon les époques, les traditions linguistiques ou orthographiques, la langue de l'auteur ou la langue cible de la traduction, les transcriptions sont extrêmement variables. En ce qui concerne le sens, deux champs sont réservés à des interprétations simplifiées (respectivement en français et en anglais) de la traduction originale. Comme pour les formes, cette simplification sert à ce que lors des recherches et des tris, des sens semblables soient affichés ensemble. Par

² Cf. <http://www.glottopedia.org/index.php/Doculect>

exemple, il n'est pas rare que le même mot désigne à la fois 'le bras' et 'la main'. Mais suivant les cas, les traductions proposées peuvent présenter des variations, même mineures, pouvant rendre plus difficile la comparaison : 'bras, main' ; 'main, bras' ; 'le bras, la main' ; 'hand, arm' ; 'arm, hand' ; etc. Une recherche sur le terme 'bras' par exemple ne renverra pas les mots dont la traduction originale est en anglais ; les définitions commençant par 'main' seront affichées loin de celles commençant par 'bras', etc. On a donc jugé utile d'ajouter, à côté de la traduction d'origine, deux champs (en français et en anglais) contenant une version relativement standardisée de cette traduction.

Bien sûr, la simplification de la traduction est beaucoup plus complexe que la simplification de la transcription, et il subsiste dans ReFLex de nombreuses irrégularités. Mais le principe même, malgré un énorme coût en temps de préparation, s'est montré très efficace.

Au mois de juin 2015, le nombre des entrées lexicales dans ReFLex a dépassé le million.

1.5. Les outils

Outre le module d'aide à la reconstruction présenté ci-dessous, ReFLex propose un module statistique et un module cartographique.

- Le module statistique permet de réaliser toutes sortes de comptages et d'afficher les résultats sous différentes formes, dont des tables de contingences pour les statistiques dites 'combinatoires' et les statistiques dites 'croisées'. Les statistiques combinatoires comptent des combinaisons à l'intérieur d'un champ unique. Typiquement, on les utilise pour explorer les combinaisons de phonèmes et/ou de tons dans les langues. Les statistiques 'croisées' permettent d'évaluer la dépendance entre le contenu de deux champs, par exemple la structure syllabique et les parties du discours. L'obtention de ces statistiques et leur exploitation scientifique sont évoqués dans ce volume (KP, réf?).

- Le module cartographique affiche le résultat d'une recherche sur une carte d'Afrique, chaque langue étant représentée par un point. Il s'agit d'un véritable module, et non d'une simple fonction, car les cartes ainsi obtenues peuvent être modifiées par l'utilisateur, par exemple pour afficher certaines informations (formes, sens, nom de langue) ou pour modifier certains paramètres (forme et couleur des points, ajout de données, etc.).

Ces deux modules peuvent être utilisés indépendamment de la base de données, ce qui n'est pas le cas du module de reconstruction.

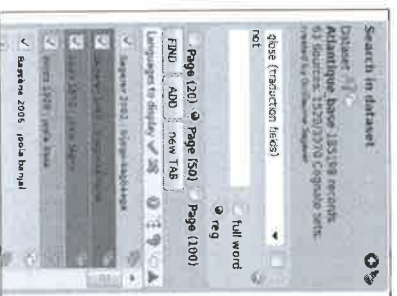
2. LE MODULE D'AIDE À LA RECONSTRUCTION

La puissance de calcul de l'ordinateur peut être d'une aide précieuse dans les tâches fastidieuses de la reconstruction phonologique et lexicale. Cependant, aucune des diverses tentatives pour remplacer complètement l'esprit humain dans cette entreprise n'ont donné de résultats convaincants. Qu'il s'agisse de rechercher automatiquement des cognats potentiels, ou d'aligner correctement les segments au sein des séries lexicales, l'algorithme n'est pas pour l'instant

en mesure de fournir des réponses correctes à tous les pièges que posent les langues naturelles. En revanche, les capacités de l'ordinateur à trier, stocker, retrouver, organiser l'information peuvent être utilisées pour soulager le linguiste et lui apporter un maximum de confort de travail. C'est dans cet esprit qu'a été conçu le module d'aide à la reconstruction de ReFLex.

D'une manière générale, l'utilisateur de ReFLex manipule des *objets*. Ces objets peuvent avoir une existence préalable (les sources, les fiches lexicales) ou être créés par l'utilisateur. Le premier objet qu'il est nécessaire de créer pour utiliser le module de reconstruction est le *dataset*. Il s'agit de l'ensemble de sources au sein duquel se fera la recherche de cognats (cet anglicisme désigne une série lexicale dont les éléments sont supposés descendre d'un étymon commun)³. L'intérêt de travailler sur un nombre réduit de sources découle de la nature même de ReFLex, qui a pour vocation à accueillir toutes les données disponibles, indépendamment de leur qualité. Dans le cadre d'un travail de reconstruction, il est important de ne sélectionner que les sources les plus dignes de confiance, lorsque le choix existe.

A partir du dataset, l'objectif est de créer des cognats, ou *cognate-sets* (voir ci-dessus). La méthode la plus simple est de lancer une recherche à l'aide de l'interface appropriée :



Par défaut, la recherche porte sur tous les champs dédiés au sens : la traduction d'origine, mais aussi les champs de traduction simplifiée en français et en anglais. Il est bien sûr possible de lancer une recherche sur un autre champ, comme il est possible de sélectionner certaines sources seulement, ou de combiner plusieurs critères de recherche. D'autres fonctionnalités (expressions régulières, exclusion) rendent l'outil de recherche extrêmement puissant.

L'affichage des résultats se fait sur la zone centrale de l'espace, sans que ne disparaisse l'interface de recherche. Par défaut, seuls quelques champs sont

³ Il importe dès à présent de signaler que beaucoup des objets de ReFLex peuvent être partagés entre plusieurs utilisateurs. Le partage d'un objet implique le partage des objets qui en dépendent, ce qui permet un véritable travail collaboratif.

Le tableau peut être trié sur chacune des colonnes, par un clic sur son en-tête. En cliquant sur le nom d'un cognate-set à gauche, on peut éditer celui-ci, et modifier l'alignement ou le nom des correspondances.

Il serait évidemment très utile d'obtenir automatiquement l'inventaire des correspondances les plus régulières, et cette option est actuellement à l'étude. Mais cette automatisation n'aura pas que des conséquences positives, dans la mesure où les séries les plus triviales seront favorisées, risquant de masquer des correspondances plus exotiques. Dans certain cas en effet, une correspondance unique peut être extrêmement convaincante et influencer sur la recherche des cognats et la découverte d'autres correspondances du même type, comme va le montrer l'exemple suivant, tiré de notre travail sur les langues atlantiques :

Dans ces langues, le mot pour 'œil' est extrêmement stable, ce qui signifie que presque toutes les langues documentées présentent une racine lexicale semblable. Voici la plupart de ces formes, dont on a isolé la consonne finale pour faire apparaître la correspondance⁶ :

LANGUE	SOURCE	'OEIL'	C#
Bijogo-kagbaaga	Segerer 2002	ne / ɲe	0
Joola fogny	Sapir 1970	ɕil	1
Joola kasa	Wintz 1909	jikil ~ jikil / ku-	1
Joola banjal	Bassène 2006	ji-cil	1
Joola kwaaray	Payne 1992	ɕkin	n
Karon	Sambou 2007	nikin	n
Bayot	Diagne 2009	sio	o
Manjaku bok	Doneux 1975	kas	s
Manjaku de Bassarel	Buis 1990	pekas	s
Pepel	Ndao 2010	kil	1
Mancagne	Trifkovic 1969	kef	f
Balante ganja	Creissels & Biaye 2015	f-git / g-git	t
Balante fraase	Wilson 2008	f-kit / k	t
Balante kentohe	Doneux & al 1984	fkit / kkit	t
Gubaher	Cobbinah 2013	sijil	1
Baynunk guñaamolo	Bao Diop 2013	si-gil / i-gil	1
Cobiana	Wilson 2008	si-gih / ji	h
Casanga	Wilson 2008	si-gir / ga / ji	r
Basari	Ferry 1991	a-ngès	s
Tanda	Wilson 2008	a-ngaz / ba	z
Bedik	Ferry 1991	gi-ngus	s
Bapen	Ferry 2006	a-ngas / ba-	s
Konyagi	Santos 1996	i-nkër / vi-nkär	r
Biafada	Wilson (p.c.)	(w)-gara / maa-g	r
Badiaranke	Meyer 2001	mase	s
Palor	d'Alton 1987	zil	1
Saafi	Mboodi 1983	has	s
Noon	Williams & Williams 1993	k'as	s
Laala	Pichl 1981	ko'as	s
Wolof	Fal, Santos & Doneux 1990	bat	t

⁶ La ligne horizontale dans le tableau sépare les langues du groupe Bak (au-dessus) des langues du groupe Nord (au-dessous).

Pular	Bah 2009	yitere	t
Peul Adamawa	Touneux 1999	yitere / gite	t
Peul Massina	Seydou 2014	yitere / gite	t
Sereer	Crétols 1973/77	ngid a...al / kid a...ak	d
Baga Mboteni	Ferry (p.c.)	cil / si-lyir	r
Nalu	Seidel 2013	cet	t
Baga fore	Golovko 2015 (p.c.)	cët / ci-ncët-ël	t

Comme on le voit, cette correspondance n'est pas triviale, puisqu'on y trouve les éléments suivants : Ø, d, h, l, h, n, r, s, j, t, z.

La série 'œil' est la plus complète. Cependant, d'autres séries réduites présentent les mêmes correspondances, ou plutôt des correspondances compatibles :

LANGUE	C#	'MORDRE'	'INTESTIN'	'TROIS'
Bijogo-kagbaaga	0			
Joola fogny	1			
Joola kasa	1			
Joola banjal	1			
Joola kwaaray	n			
Karon	n			
Bayot	o			
Manjaku bok	s	pas		
Manjaku de Bassarel	s	pes		
Pepel	1	pul		
Mancagne	f	pef		
Balante ganja	t	m-mbüté		
Balante fraase	t	b-mbuté		
Balante kentohe	t	kambute		
Gubaher	1	bungal		lall
Baynunk guñaamolo	1	ɲah		lall
Cobiana	h	ɲar		
Casanga	r	a-ŷas		tās
Basari	s	u-ɲās		tās
Tanda	z			tās
Bedik	s	i-ɲër		tær
Bapen	r	ɲar-	bu-bur	
Konyagi	r	ɲas		
Biafada	r			
Badiaranke	s			
Palor	1			
Saafi	s			
Noon	s			
Laala	s			
Wolof	t	mat	buitt	tat
Pular	t	ɲat		tat
Peul Adamawa	t	ɲat		tat
Peul Massina	t	ɲat		tad
Sereer	d	ɲerà	fud o...ol	
Baga Mboteni	r		nio-bürük	
Nalu	t	ɲat		paat
Baga fore	t			tet

Ce qu'on sait des correspondances possibles grâce à la série 'œil' va permettre de rechercher des cognats pour les séries incomplètes, et/ou d'éliminer des faux cognats. Ainsi, pour la série 'mordre', j'avais dans un premier temps intégré une racine *ɲar* pour les langues du groupe joola, avec le sens 'prendre'. Le changement sémantique supposé 'prendre' <> 'mordre' paraissait plausible. Mais la consonne finale ne s'accordait pas avec celle prévue par la série 'œil'. Celle-ci aurait dû être *l*. La série 'œil' m'a donc incité à rechercher une forme *ɲal* dans le groupe joola. Cette forme existe en joola foggy et en banjal-guslay avec le sens de 'mâchoire' et en joola kasa avec le sens de 'joue'. La variation sémantique 'mordre' <> 'mâchoire' (~joue) est assez plausible elle aussi. On a ainsi découvert une racine vraisemblablement proto-atlantique, avec une probable innovation sémantique en joola.

Le module de reconstruction intègre un grand nombre de fonctionnalités, parfois expérimentales, dont certaines sont encore en cours de développement à l'heure où ces lignes sont écrites. En voici quelques-unes :

- Au moment de nommer une correspondance après avoir procédé à l'alignement des cognats, le bouton 'search' (visible sur la figure 3 ci-dessus p. 207) doit permettre de trouver, parmi les alignements déjà enregistrés, ceux qui sont plus ou moins compatibles avec l'alignement en cours.
- Plusieurs alignements différents peuvent être proposés sur un même cognate-set, ce qui permet de tester différentes hypothèses.
- Un cognate-set est par défaut conçu comme une étape vers la reconstruction, via l'alignement. Mais d'autres options sont possibles, qu'il serait trop long de détailler ici.
- Un petit module permet de proposer aux autres utilisateurs d'un dataset partagé de modifier un cognate-set, par l'ajout ou la suppression de fiches. Ces propositions de modification peuvent être acceptées ou non par le destinataire, et/ou donner lieu à des discussions.
- Pour faciliter le travail d'alignement sur de grandes quantités de cognate-sets, leur affichage peut être limité à ceux dont l'alignement n'est pas encore enregistré.

CONCLUSION

La reconstruction phonologique et lexicale est un domaine ou la linguistique africaine est souvent considérée comme déficiente. Pendant longtemps, il était d'usage d'invoquer le manque de données, ou la difficulté d'accès aux données. Le projet ReFlex vise justement à mettre un terme à ces obstacles. Aujourd'hui les données sont très nombreuses : on peut estimer que l'ensemble des dictionnaires et lexiques publiés concernant les langues africaines contient au moins 10 millions de mots, pour au moins 1000 à 1500 langues. ReFlex permet d'accéder très facilement à une partie non négligeable de ce corpus, et l'ensemble des données accessibles est appelé à continuer à croître. Même si de nombreux groupes de langues sont encore insuffisamment représentés pour pouvoir envisager une reconstruction sérieuse (notamment dans les groupes

Mande, Gur, Kwa, Kru), les possibilités et le confort d'utilisation proposés par ReFlex devraient inciter les linguistes intéressés par la reconstruction à y intégrer davantage de données. Déjà, le travail sur les langues Atlantiques et en bonne voie.

Par ailleurs, l'équipe de DDL à Lyon est en train d'élaborer une base de données sur les langues amérindiennes, qui a la même structure que ReFlex et dispose donc des mêmes outils. D'autres initiatives similaires, concernant d'autres régions du monde, sont possibles et bienvenues.

BIBLIOGRAPHIE

- Barry A., 1987, *The joola languages: subgrouping and reconstruction*, London, School of Oriental and African Studies (SOAS), University of London (PhD thesis).
- Carlton E. M. & Rand S. R., 1994, Compilation de listes de mots Swadesh modifiées recueillies parmi les langues du groupe Diola, sud du Sénégal, *Journal of West African languages* 24-2, p. 85-119.
- Fresco E. M. 1970, *Topics in Yoruba dialect phonology*, Los Angeles, African Studies Center & Department of Linguistics, University of California (UCLA).
- Sambou P., 2007, *Morphosyntaxe du joola karon*, Dakar, Université Cheikh Anta Diop (Thèse de doctorat de 3^e cycle).
- Snider K. L. 1989, *North-Guang comparative wordlist: Chumburung, Krachi, Nawuri, Gichode, Goniya*, Legon, Institute of African Studies, University of Ghana.
- Wilson W. A. A., 2007, *Guinea Languages of the Atlantic Group*, Frankfurt, Peter Lang.