

Informatique et terminologie : revue des technologies nouvelles

Pierre Auger

Volume 34, numéro 3, septembre 1989

1. Actes du Colloque Les terminologies spécialisées : Approches quantitative et logico-sémantique et 2. Actes du Colloque Terminologie et Industries de la langue

URI : <https://id.erudit.org/iderudit/001923ar>

DOI : <https://doi.org/10.7202/001923ar>

[Aller au sommaire du numéro](#)

Éditeur(s)

Les Presses de l'Université de Montréal

ISSN

0026-0452 (imprimé)

1492-1421 (numérique)

[Découvrir la revue](#)

Citer cet article

Auger, P. (1989). Informatique et terminologie : revue des technologies nouvelles. *Meta*, 34(3), 485–492. <https://doi.org/10.7202/001923ar>

INFORMATIQUE ET TERMINOLOGIE : REVUE DES TECHNOLOGIES NOUVELLES

PIERRE AUGER
Université Laval, Québec, Canada

On parle beaucoup, ces dernières années du rôle de plus en plus important que joue l'informatique dans le développement des sciences humaines. Certaines disciplines, plus que d'autres, ont profité de ce moyen puissant de traitement, je veux parler ici de tout le domaine des sciences du langage et plus particulièrement des applications qui en dérivent et qui font appel à la manipulation importante de données. La terminographie est une de ces applications qui a bénéficié du traitement automatique des données, cette rencontre de la terminologie appliquée et de l'informatique a créé un champ nouveau de travail qu'on a dénommé la terminotique. Pour cet exposé, nous nous situerons du côté de l'informatique linguistique contemporaine et des produits «industriels» qui sont en amont ou en dérivent pour nous approcher du concept plus englobant d'industries de la langue (IDLL) à travers l'exploration du concept de terminotique.

La terminotique peut ainsi être considérée comme un concept de recherche particulier de l'informatique linguistique appliquée, celui du traitement automatique du terme, mais surtout un concept pragmatique étroitement relié à ce qu'on appelle aujourd'hui l'industrialisation de la langue, un concept générateur de produits performants de traitement de l'écrit dont les applications sont faciles à dégager :

- ◆ constitution automatique de bases de données terminologiques,
- ◆ élaboration automatique de dictionnaires terminologiques,
- ◆ préparation de dictionnaires spécialisés pour assister les systèmes de T.A.O.,
- ◆ mise au point de bases de connaissances pour les systèmes-experts «intelligents», etc.,
- ◆ poste de travail automatisé «intégré» pour le traducteur
- ◆ élaboration de thésaurus documentaires multilingues en langue naturelle,
- ◆ accès multilingue aux banques de données,
- ◆ génération automatique de textes.

Comme introduction à notre propos, il convient de dire que le concept de terminotique s'appuie avant tout sur les progrès accomplis par l'informatique d'orientation textuelle depuis dix ou quinze ans. Conçus à l'origine pour le traitement de données numériques, les ordinateurs se sont peu à peu adaptés au traitement de données textuelles (les caractères). Une deuxième constatation s'impose également qui nous fait observer que l'on doit beaucoup à la sociabilisation de l'informatique et à sa portabilité, qui ont permis à des linguistes ou des langagiers (comme on les appelle maintenant au Québec) d'accéder à l'informatique et d'utiliser directement ce moyen pour leurs recherches et même de l'adapter plus efficacement à leurs besoins. Ces deux constatations constituent en quelque sorte le fondement historique du vaste domaine de l'informatique langagière et de ses applications majeures.

Il faut également mentionner les besoins sans cesse croissants en terminologie vive engendrés par le flux informationnel spécialisé (techno-scientifique) des années 80.

Mentionnons à titre d'exemple le marché actuel de la traduction, grande consommatrice de terminologies vives. Ce marché mondial de la traduction représente aujourd'hui 150 millions pages/an avec un contingent de 175 000 traducteurs. Le marché européen de la traduction (CCE), pour sa part, s'étend sur douze pays et neuf langues, ce qui représente 72 couples de langue à traduire répartis sur 1,5 millions de pages/an pour un effectif de plus de 2 000 traducteurs. Si l'on analyse cette fois les domaines généraux faisant l'objet de cette intense activité traductionnelle, on aura la ventilation approximative suivante : 70% pour la traduction de documents administratifs, 30% pour la traduction de documents techniques et 4% seulement pour la traduction littéraire. Ces chiffres donnés par J.-F. Dégremont lors du Congrès de l'Aupelf, à New-Delhi, parlent d'eux-mêmes de ce flot terminologique intense amené par les communications modernes et qui va s'intensifiant avec les années.

Pour conclure cette introduction, disons également que la terminologie, comme de nombreuses disciplines reposant sur le traitement de l'information écrite a évolué naturellement dans le sens d'un rapprochement constant avec l'informatique et ses moyens. En effet, la terminographie (comme la lexicographie) repose sur un traitement extensif et intensif de l'information textuelle : *extensif* parce qu'elle traite du texte sous toutes ses formes, prenons pour exemple la diversité des informations nécessaires à l'accomplissement d'un travail terminologique et les tâches de traitement qui s'y rattachent ; *intensif* : songeons seulement au volume considérable de cette documentation et aux lourdes tâches, souvent répétitives, qui attendent le terminologue dans sa recherche, on se rapproche ici sensiblement d'un idéal méthodologique d'exhaustivité que permet seul l'automate. (Ceci a légitimé l'approche que nous avons privilégiée à l'Université Laval pour la formation des traducteurs et terminologues, orientation vers les besoins actuels et futurs de l'industrie.)

Si l'on examine de près le processus méthodologique lié à la confection d'un dictionnaire terminologique du début des travaux jusqu'à sa mise en marché, quatre grandes catégories de tâches peuvent être distinguées :

- 1) l'accomplissement de tâches préliminaires qui sont avant tout d'ordre documentaire (exploration du domaine, sélection et enregistrement des sources) ;
- 2) l'accomplissement de tâches proprement terminographiques liées au traitement d'écrits spécialisés (mise en forme et traitement d'un corpus de textes) ;
- 3) le stockage et le traitement des données recueillies (structuration de bases de données et utilisation de bases de données) et
- 4) l'édition et la diffusion du produit final sous diverses formes.

Il est maintenant facile d'imaginer que chacune de ces quatre catégories de tâches traditionnellement conduites manuellement peut, séparément, faire l'objet d'un traitement automatique, et cela à des degrés divers d'automatisation. Une projection vers un avenir qui n'est peut-être pas très lointain permet d'imaginer sans doute l'automatisation complète de l'ensemble de la chaîne de traitement terminographique depuis les premières recherches documentaires jusqu'à l'obtention du produit final, le dictionnaire terminologique, quelque soit sa forme (imprimée ou électronique). Voyons maintenant ce que peut apporter l'informatique comme assistance aux tâches terminographiques.

A. L'INFORMATIQUE COMME SOUTIEN AUX TÂCHES TERMINOGRAPHIQUES

1. *La phase documentaire* : Le recours à la télématique pour la téléconsultation et le téléchargement à partir de banques de données bibliographiques, textuelles et terminologiques, peut raccourcir de beaucoup l'étape de l'établissement du corpus de travail (sources, documents de référence, éléments de connaissance, etc.). On observe souvent que ces outils disponibles sont rarement exploités de façon intelligente ou optimale.

2. *La mise en forme et le traitement du corpus de textes spécialisés (sur deux langues, par exemple) :*

- ◆ le téléchargement ou saisie optique ou manuelle pour constituer le corpus écrit ;
- ◆ la mise en forme du corpus saisi par un logiciel de traitement de texte ;
- ◆ la lecture et le traitement du fichier-texte par un logiciel de découpage de mots/termes et de termes complexes qui permet l'établissement automatique ou semi-automatique (mode interactif) de la nomenclature ;
- ◆ la lecture et le traitement par un logiciel d'indexation (découpage des contextes, comptage des formes, analyse de fréquence, repérage des descripteurs) ;
- ◆ le découpage simultané des contextes ;
- ◆ l'analyse sémantique des contextes en regard des descripteurs pour la rédaction assistée des définitions.

3. *L'élaboration d'une base de données :*

- ◆ la structuration et le stockage de l'information avec un s.g.b.d (original ou commercial) ;
- ◆ le traitement de l'information (ajouts, retraits, modifications, tris de données structurées) ;
- ◆ la fusion des bases de données, anglaise et française, par exemple, à partir d'un thésaurus) ;
- ◆ l'édition de la base de et la mise en forme des données d'un fichier-texte.

4. *L'édition de la base de données :*

- ◆ le transfert du fichier-texte sur un système d'édition pour arriver à un prêt-à-publier électronique ;
- ◆ la publication sous différents supports (disquettes, CD-ROM, vidéo-disques, disques WORM, imprimés divers).

La liste un peu longue qui vient d'être livrée, sans être très affinée du point de vue du traitement informatique, suggère qu'à toutes les étapes du travail terminologique, une forme d'automatisation, variable en intensité, est possible en recourant ou non à l'interface humain. L'idéal serait la réunion en un progiciel facilement utilisable d'un ensemble de programmes-outils pour l'automatisation de chacune des phases de travail.

Le module principal de tout système de terminotique est le système de gestion de bases de données (sgbd). Plusieurs alternatives se présentent quant au choix d'un logiciel de sgbd. Le terminologue ou le terminoticien peut d'abord avoir recours à un logiciel commercial standard de sgbd, tels Dbase3+ ou 4, Foxbase, Edibase, Knowledgeman ou ZIM (ZANTHE) qui ne sont pas construits spécifiquement pour le traitement linguistique, mais qu'il pourra adapter au traitement de données terminologiques en exploitant au maximum les caractéristiques de ces logiciels «prêts-à-porter» qui sont presque exclusivement de type relationnel (De Schaetzen, 1987a, 1987b, 1987c). Caroline de Schaetzen en décrivant les possibilités de divers systèmes expose également leurs limites : langage procédural syntaxiquement difficile, importance des mémoires de masse nécessaires, manque de souplesse pour le traitement de certaines langues, africaines, entre autres, mais surtout leur relative «bêtise» qui rend difficile, voire impossible tout prise en compte sémantique (ce qui serait différent en utilisant des langages de programmation recourant à l'I.A. comme LISP ou PROLOG pour créer des systèmes experts comme GURU). On est encore loin de la «station de travail lexicographique et linguistique dans laquelle le linguiste et le lexicographe peuvent interagir avec les différentes sources de connaissances : interagir avec les dictionnaires qui existaient [sic], les corpus de référence, etc.» (ZAMPOLLI, 1988).

Comme solution transitoire entre la base de données «bête» et le système expert basé sur un cumul de connaissances, des progiciels de terminotique ont été développés ces dernières années qui «roulent» sur micro-ordinateur (PC compatibles) et qui permettent de constituer des bases de données terminologiques relationnelles et de les éditer et qui, de plus, sont adaptés au traitement terminologique systématique en tenant compte de la structuration notionnelle du domaine traité. Signalons, parmi les plus connus, le système BATEM de J. Baudot en développement à l'Université de Montréal et le système MICROCEZEAU de Cézeauterm (Université de Clermont-Ferrand) qui fonctionnent sur la norme IBM-PC. Le dernier est disponible commercialement en plusieurs versions, la plus récente étant programmée en langage DBASE3+. Notons toutefois que ce dernier progiciel est avant tout un sgbd adapté au traitement terminologique et qu'il ne permet pas l'automatisation complète de la chaîne de travail terminologique. D'autres produits du même genre sont apparus sur le marché récemment; ils ont en commun d'offrir les mêmes limites quant au degré d'automatisation offert. Il faut ici encore parler d'une autre catégorie de produits apparentés (par ex. des gestionnaires de mini-bases de données terminologiques multilingues) qui sont des outils plutôt orientés vers l'assistance à la traduction, comme TERMEX, INK TEXTTOOLS, ALPS SELECTERM et AUTOTERM ou ABC-WORD, qui permettent à un traducteur de constituer ses propres dictionnaires terminologiques automatisés et qui étant chargeables dans la mémoire résidente d'un micro-ordinateur peuvent être interrogés directement par le traducteur qui utilise un traitement de texte. Il va sans dire que cette dernière catégorie de logiciels, de par ses limites mêmes, ne peut être classée parmi les progiciels de terminotique.

Enfin, il reste encore à explorer les applications possibles à tirer de bases de données libres (non structurées) constituées par des «hypertextes» gérés par des gestionnaires «en piles» de type Hypercard (forme plus évoluée des logiciels de recherche textuelle comme WORDCRUNCHER, GOFERS, IZE, Zy-Index).

B. LES PROBLÈMES RENCONTRÉS

1. *Problèmes généraux*: Si les outils informatiques d'analyse et de traitement de l'écrit (spécialisé, dans le cas qui nous intéresse) existent aujourd'hui (p.ex. Ink Tools, Wordcruncher, Alps Tools qui sont des outils commercialisés, etc.), ou ils ne sont pas assemblés en progiciels complets ou ils sont trop difficiles d'utilisation pour des non-initiés (songeons aux lemmatiseurs en particulier et aux divers logiciels de découpage ou d'indexation de mots/termes), de plus, ils nécessitent encore, pour une bonne part, une intervention humaine constante pour valider les choix proposés par la machine. Prenons par exemple les découpeurs de termes complexes qui ne peuvent fonctionner en mode complètement autonome (par ex. sans intervention humaine) et qui reposent sur l'élaboration et l'enrichissement constant d'un dictionnaire de référence où les choix sont fournis à la machine par l'opérateur. Ces outils sont également diversement performants quant à leur aptitude à reconnaître, traiter ou analyser la chaîne de caractères du français ou encore sa syntaxe. Problème plus général encore, ces progiciels construits autour de sgbd ne fonctionnent pas en langage naturel et nécessitent pour l'opérateur l'utilisation d'une syntaxe souvent très complexe pour effectuer des tâches spéciales, d'où la nécessité de recourir à l'élaboration de menus fastidieux et de routines d'automatisation difficiles à générer.

Comme remède à ces problèmes généraux, il faut d'abord compter sur l'«industrialisation du français» et le développement de sa capacité de dialoguer avec les ordinateurs (c'est déjà fait pour la langue anglaise) et par ricochet, le développement d'une informatique en langue française naturelle. Songeons, à titre d'exemple, aux limites d'accès à l'information dans les banques de terminologie qui utilisent un langage d'interrogation

procédural très éloigné du langage naturel. Ces banques sont toujours des systèmes d'information gigantesques sous-exploités fonctionnant avec de vieux moteurs. Leurs limites sont imposées par leur langage d'interrogation vétuste, par le peu d'«intelligence» qu'ils ont à répondre aux besoins des utilisateurs. L'avenir en ce domaine réside très certainement dans le développement de logiciels reposant sur l'intelligence artificielle fondant des systèmes experts langagiers capables de manipuler efficacement le langage humain. L'avenir des systèmes de T.A.O., à titre d'exemple, se situe très exactement dans cette problématique de développement.

En examinant les aspects logiciels en micro-informatique, la situation n'est guère plus brillante :

- ◆ adéquation de l'équipement (matériel) : exigences de mémoires de masse et de puissance de l'unité centrale de traitement qui croissent sans cesse ;
- ◆ compatibilité des équipements : songeons seulement à la possibilité d'intégration de ces systèmes dans un contexte bureautique ;
- ◆ convivialité et ergonomie des systèmes existants ;
- ◆ limites des logiciels commerciaux standards disponibles.

2. *Problèmes spécifiques :*

Le problème le plus sérieux, à notre avis, dans la perspective du développement de la terminotique est la disponibilité des textes spécialisés en version électronique (lisibles par la machine).

Les systèmes de lecture optique pour la saisie automatique des textes écrits ne comblent que de façon très partielle cette lacune. Les systèmes de saisie optique de caractères se situant dans une gamme de prix abordable sont souvent limités quant à la variété de caractères qu'ils peuvent reconnaître ou sont paresseux pour «apprendre» ou sont peu fiables quant au taux d'erreur qu'ils produisent. De toute façon, les textes saisis par lecture optique nécessitent dans les phases ultérieures de traitement, de nombreuses manipulations (suppressions de toutes sortes, formatage, encodage) difficiles à exécuter par des novices et qui sont coûteuses en temps. C'est ici que des bases de données textuelles réunissant des textes spécialisés récents en langue française et couvrant un large éventail de domaines deviendraient des outils de travail extraordinairement puissants pour le terminoticien [à l'exemple de l'anglais avec le BROWN corpus (1 million de mots) ou The Birmingham Collection of British English (22 millions de mots)]. On pourrait alors songer au téléchargement à partir de ces banques et même arriver à constituer des corpus de dépouillement équilibrés et représentatifs. L'exploitation des milliers de bases de données informatives spécialisées existantes avec des logiciels appropriés de recherche textuelle est une urgence à ne pas négliger pour constituer rapidement des fonds terminologiques «exhaustifs». Songeons également aux bandes magnétiques des journaux et revues générales et spécialisées qui pourraient être exploitées de la même façon.

D'un autre point de vue, on pourrait songer à utiliser de la même façon les dictionnaires, tant généraux que spécialisés, sous leur forme électronique (tous les dictionnaires importants existent sous cette forme avant d'être publiés). Ces «hypertextes» traités, par exemple, par des logiciels d'indexation pourraient se révéler être de précieux aides pour le terminologue (recherches de contextes, rédaction de définitions, etc.). À plus forte raison encore les dictionnaires et les encyclopédies existant sous forme de CD-ROM (malheureusement presque exclusivement en langue anglaise jusqu'à maintenant, comme *The Oxford English Dictionary* ou le Bookshelf de Microsoft qui offre sur un seul disque compact toute une gamme d'ouvrages de référence dont *The American Heritage Dictionary*). Il faudra probablement attendre encore quelques années pour voir apparaître de tels outils en français (p. ex. Le Trésor de la langue française de Nancy). Enfin, il y a

les banques de terminologie qui peuvent jouer le même rôle, étant pour la plupart dotées des facilités pour constituer des fichiers électroniques à partir du matériel extrait lors de l'interrogation.

C. PROSPECTIVES ET VOIES D'AVENIR

Faisons maintenant table rase des problèmes que nous venons d'identifier et considérons qu'ils sont d'ores et déjà réglés. Dans ce scénario futuriste, le terminologue, à partir de son poste de travail automatisé a accès à de gigantesques bases de données (ou de connaissances) textuelles, il extrait par téléchargement les éléments de son corpus de ces bases de données, le fait dépouiller automatiquement sans avoir eu à saisir le texte manuellement au préalable, établit automatiquement sa nomenclature de travail, fait découper les termes-entrées (le plus souvent complexes pour une langue comme le français) par la machine, en extrait des descripteurs sémantiques qui serviront ultérieurement à rédiger des définitions en mode assisté, classe, trie, fusionne les bases de données et les édite avec un minimum d'intervention de sa part. Son poste de travail multi-tâches lui permet tout en poursuivant son travail d'accéder instantanément à des banques de données terminologiques ou documentaires, d'échanger avec ces systèmes et de télécharger de l'information dans son propre système de terminotique qu'elle soit textuelle ou graphique (une illustration, par exemple), d'utiliser une abondante documentation sur CD-ROM, enfin, doté de fonctions bureautiques avancées «intelligentes», ce poste de travail lui permet de contrôler lui-même et en tout temps l'élaboration de son produit et de le mener à terme dans les meilleures conditions. Nous avons là, en fait, toutes les composantes d'un système expert de terminotique capable de prendre lui-même certaines décisions et de les faire corroborer au besoin par le terminologue qui se réserve alors le rôle valorisant de contrôler les faits et gestes de son automate. De plus, les problèmes de mise à jour des dictionnaires terminologiques sont aplanis dans un tel scénario. Cet horizon idyllique est peut-être plus près de nous que nous ne l'imaginons ; attendons ce que nous réservent les prochaines années en ce domaine.

Il faudra certainement longtemps compter encore sur les progrès de la recherche en informatique linguistique et sur les travaux systématiques de centres comme le LADL (lexique et syntaxe du français), ATO à l'UQUAM (dépouillement automatique de termes et termes complexes, cf. projet avec OLF) et bien d'autres encore, avant de voir résolus les problèmes que nous venons d'évoquer. Mentionnons également des centres d'application comme le CRITT (au Québec) ou la la DIST en France (Délégation de l'information scientifique et technique) dont la mission est de soutenir le développement de projets et d'applications d'automatisation linguistique. Pour la terminotique spécifiquement, une coordination internationale paraît nécessaire pour éviter les dédoublements et orienter les efforts de développement selon des axes coopératifs. Les travaux du groupe ISO/TC 37/SC 3 N 37 en «Computer aids in terminology (working draft)» se situent dans cette voie mais pas toujours selon une approche réaliste d'unification des méthodes de travail (cf. poids de certains groupes linguistiques dans le Comité 37 d'ISO).

Pour donner un aperçu des nouveaux outils de traitement du langage que l'informatique met déjà à notre disposition et qui seront perfectionnés encore dans les prochaines années pour servir le terminologue, mentionnons :

- ◆ les systèmes intelligents de reconnaissance de caractères ;
- ◆ les analyseurs syntaxiques (parsers) et les lemmatiseurs capables de décortiquer un texte en unités et de les classer grammaticalement ;
- ◆ les découpeurs de mots / termes et de termes complexes ;
- ◆ les analyseurs sémantiques ;
- ◆ les logiciels d'indexation ;

- ◆ les logiciels d'interrogation en langue naturelle ;
- ◆ les nouvelles technologies de diffusion de l'information terminologique (CD-ROM, vidéo-disques, disques WORM, etc.) ;
- ◆ et enfin, des banques de terminologie « intelligentes ».

Ce que nous avons voulu mettre en évidence dans le développement que nous venons de faire, ce sont les possibilités énormes que l'informatique offre désormais au lexicographe et au terminologue pour qui les problèmes de temps, d'efficacité et de rentabilité sont toujours présents. Les responsables de la réédition du grand dictionnaire anglais *The Oxford Dictionary of English* (entreprise logée au Canada, à Waterloo) ont affirmé récemment qu'une entreprise de cette ampleur (p. ex. la réédition du ODE) est désormais envisageable dans un délai d'une année, alors qu'il y a quelques années encore, une telle tâche conduite manuellement aurait duré quelque soixante années, avec moins de fiabilité et de perfection. Cette relation d'efficacité entre l'homme, la langue et la machine porte aujourd'hui un nom et c'est l'**industrialisation de la langue** et les technologies qu'elle met en œuvre se regroupent sous l'étiquette **industries de la langue**. Disons enfin que l'atteinte de cet objectif de développement de produits industriels en langue française est primordial pour le maintien du statut international du français comme grande langue véhiculaire de la science et des techniques.

En terminant, j'aimerais insister sur la nécessité qu'il y a pour les langagiers de notre génération de se frotter à toute cette nouvelle technologie d'industrialisation de la langue. Les développements en ce domaine nous concernent tous que nous soyons des concepteurs ou des utilisateurs de l'informatique. Nous avons tenté ici d'illustrer les bénéfices qu'on peut tirer d'une intégration forte des activités langagières à l'informatique avec les moyens nouveaux qu'elle met à notre disposition. Conscient des besoins urgents des professions langagières en ce domaine, le Département de langues et linguistique de l'Université Laval a choisi de former ses étudiants linguistes, terminologues et traducteurs selon cette orientation. Les cursus prévoient pour chaque programme une initiation poussée à l'informatique de texte et à l'informatique linguistique dans certains cas toujours selon une approche très pratique. Ce virage de notre institution a été rendu possible grâce à l'accord IBM/Laval conclu en 1985. Cette orientation sera maintenue et même intensifiée au-delà de la fin de l'accord prévu pour le mois de décembre 1988. Notre objectif ultime est de former de futurs praticiens langagiers prêts pour le marché du travail et capables de manipuler l'informatique pour leurs besoins professionnels. Il s'agit également d'une orientation fondamentale nouvelle de la recherche en linguistique à Laval entérinée par le CIRB en 1987 qui entend bien contribuer au développement des Industries de la langue au Québec tout en rentabilisant la recherche universitaire.

RÉFÉRENCES

- ABBOU, André, MEYER, Thierry, LEFAUCHER, Isabelle (1987) : *Les industries de la langue. Les applications industrielles du traitement de la langue par les machines*, Paris, Éd. Daicadif, 400 p.
- Actas de la exposicion de linguistica informatica y de terminologia cientifico-tecnica* (23-38 février 1987, Madrid), Paris, Union latine, 1988.
- AUGER, Pierre (1988) : « L'industrialisation de la langue française et son maintien comme grande langue véhiculaire de la science et de la technique ». Texte d'une conférence présentée lors du 56^e Congrès de l'ACFAS, Moncton, 9-13 mai 1988, paru dans *Les industries de la langue : au confluent de la linguistique et de l'informatique*, publication K-9, Québec, CIRB, pp. 3-13.
- AUGER, Pierre (1988) : « La terminotique, un volet particulier de l'informatique langagière », texte d'une conférence présentée au Congrès annuel 1988 de la Société des traducteurs du Québec, Montréal, 6 juin 1988, (à paraître).
- AUGER, Pierre (1988) : « Le travail du terminologue amélioré et simplifié », *CIRCUIT*, septembre 1988, pp. 12-13.

- AUGER, Pierre (1988) : «Recueil de notes et textes choisis dans le cadre du cours TRD-15526. Terminologie de/et l'informatique», Département de langues et linguistique, Faculté des Lettres, Université Laval, septembre 1988, 317 p.
- AUGER, Pierre (1988) : «La terminotique et les industries de la langue», *Meta* 34-3, septembre 1989.
- BUREAU VAN DIJK (1988) : «Journées d'études 1992 : Marché unique — Marché multilingue, les outils du traitement des langues» (Paris, 15 juin 1988).
- DEGREMONT, Jean-François (1986) : «Au croisement de l'informatique et de la linguistique, les industries de la langue, Enjeux pour l'Europe, Actes du colloque de Tours, Tours, 28 fév. — 1^{er} mars 1986, *Encrages*, n° 16, nov. 1986, pp. 22-46.
- DE SCHAETZEN, Caroline (1987a) : «S.g.b.d. et terminologie», *Le linguiste (De Taalkundige)*, v. 33-3.
- DE SCHAETZEN, Caroline (1987b) : «Terminologie en langues africaines et gestionnaires de données textuelles», *Le langage et l'homme*, v. 22, fasc. 2, pp. 172-174.
- DE SCHAETZEN, C. et MEERT, O. (1987c) : «Un outil d'aide à la création et à la gestion de bases de données terminologiques pour les langues africaines» *Le langage et l'homme*, v. 22, fasc. 2, pp. 157-165.
- PARADIS, Claude, AUGER, Pierre (1987) : «La terminotique ou la terminologie à l'ère de l'informatique», *Meta*, vol. 32-2, juin 1987, pp. 102-110.
- Terminogramme* (1988) : «Terminologie & Informatique», janvier 1988, n° 46.
- ZAMPOLLI, A. (1988) : «La consultation de bases de données en langage naturel», *Actas de la exposicion de linguistica informatica y de terminologia cientifico-tecnica* (23-38 février 1987, Madrid), Paris, Union latine.