

VOLLNHALLS, Otto (1992) : *A Multilingual Dictionary of Artificial Intelligence*, London & New York, Routledge, 423 p.

Henri Béjoint et Philippe Thoiron

Volume 38, numéro 4, décembre 1993

Le *Je* du traducteur
The *I* of the Translator

URI : <https://id.erudit.org/iderudit/002026ar>

DOI : <https://doi.org/10.7202/002026ar>

[Aller au sommaire du numéro](#)

Éditeur(s)

Les Presses de l'Université de Montréal

ISSN

0026-0452 (imprimé)

[Découvrir la revue](#)

Citer ce compte rendu

Béjoint, H. & Thoiron, P. (1993). Compte rendu de [VOLLNHALLS, Otto (1992) : *A Multilingual Dictionary of Artificial Intelligence*, London & New York, Routledge, 423 p.] *Meta*, 38(4), 719–721. <https://doi.org/10.7202/002026ar>

Comptes rendus

■ VOLLNHALLS, Otto (1992): *A Multilingual Dictionary of Artificial Intelligence*, London & New York, Routledge, 423 p.

L'ouvrage (MDAI) signé Otto Vollnhalls (un nom bien connu dans le domaine de la lexicographie spécialisée) regroupe, dans un fort volume relié, une macrostructure principale de 3609 vedettes anglaises classées par ordre alphabétique et numérotées, et quatre index (allemand, français, espagnol, italien) dont les entrées renvoient aux articles par l'intermédiaire des entrées anglaises.

Sous-domaines couverts

Une liste des seize sous-domaines couverts par l'ouvrage est donnée (p. V); il s'agit de: «*expert systems, games, image processing, knowledge, learning, linguistics, logic, mt (machine translation), maths, neural nets, nlp (natural language processing), programming, recognition, robotics, synthesis, theory.*»

L'affection de certains termes à ces sous-domaines est parfois étrange. En voici quelques exemples (les noms de sous-domaines sont en CAPITALES):

■ LISP (*list processing language*) est affecté à PROGRAMMING, alors que PROLOG (*an AI programming language*) n'est affecté à aucun sous-domaine.

■ Le terme *lexical analysis* est affecté au sous-domaine NLP, alors que *lexicon* est affecté à KNOWLEDGE.

■ Les termes concernant l'étude des sons linguistiques sont éparpillés dans divers sous-domaines. Ainsi: *allophone (variant of the same phoneme)* à RECOGNITION, *phoneme* à LINGUISTICS, *phonemic data* à RECOGNITION, *phonetic transcription* à LINGUISTICS, *phonology* à LINGUISTICS, etc.

■ On peut admettre que *breadth first* soit affecté à LOGIC et *breadth first search* à EXPERT SYSTEMS si l'on distingue entre la nature d'une procédure et sa mise en œuvre. Mais comment expliquer alors que *brute force method*, qui devrait être étiqueté de la même manière que *breadth first search*, soit classé dans LOGIC, alors que l'explication qui l'accompagne (*in problem solving*) indique que pour l'auteur il s'agit bien d'une procédure? De même, il semble que l'attribution de *breadth of search* à LOGIC soit peu cohérente avec les autres articles.

■ Il n'y a pas de raisons d'affecter à des sous-domaines différents *Bayes theorem* (LOGIC) et *De Morgan's law* (THEORY). On peut également se demander pourquoi *Laplacian equation* est étiqueté MATHS¹.

■ On trouve *depth first*, que l'auteur décrit comme antonyme de *breadth first*, dans EXPERT SYSTEMS alors que *breadth first* est classé dans LOGIC. De la même manière *depth*, muni de la glose (*of a search*), est étiqueté EXPERT SYSTEMS.

Sur tous ces problèmes d'attribution à des sous-domaines, le lecteur aurait apprécié un mot d'explication dans la préface. Mais il est vrai que, dans une discipline comme

l'intelligence artificielle, la clarté et la cohérence relèvent de la gageure. Ajoutons que la décision de l'auteur de recourir à un seul sous-domaine pour un terme donné, si elle simplifie la consultation de l'ouvrage, a forcément conduit à certains choix déchirants à cause des exclusions qu'ils entraînent.

Choix des entrées

Le choix des entrées, explique l'auteur, résulte d'un vaste programme de lectures dans les cinq langues. Il est dommage que ce corpus ne soit pas détaillé, et que l'ouvrage ne comprenne même pas une bibliographie ; selon l'auteur (p. V), elle aurait été inutile tant les écrits sur le sujet sont nombreux depuis quelques années.

Si l'on compare la macrostructure de MDAI à celle de Otman (1991), on s'aperçoit que le recouvrement est assez faible. Ainsi, les entrées suivantes sont dans Otman et pas dans MDAI, qui contient pourtant davantage de termes : *actor language, actor-based language, adjacent arc, admissibility, admissible algorithm, algorithmic complexity, algorithmic parallelism, alpha-beta procedure, alpha-beta pruning, analytical knowledge, anaphora, anti demon, antidemon rule, applicative language, applicative programming, approximative knowledge, arborescence, array processing, assertional knowledge, associative access, associative redaction, associative writing, asynchronous processing, authoring system, automated translation, automatic writing recognition system*, etc. Sans doute ceci est-il dû au fait que l'intelligence artificielle est un domaine en pleine expansion, auquel contribuent des spécialistes d'horizons très divers.

Les renvois entre termes connexes sont assez nombreux et bien utiles. On sait que le nombre et la richesse de ces renvois sont différemment appréciés selon les lecteurs, et il n'est pas de critique d'ouvrage lexicographique qui n'y aille de son couplet. Le nôtre sera bref : n'aurait-il pas été utile de montrer systématiquement les liens conceptuels, surtout là où la morphologie des termes n'invite pas le lecteur à les supposer ? Ainsi, on peut imaginer que le lecteur fera le rapprochement entre *forward chaining* et *backward chaining*, et peut-être même entre *data-driven* et *goal-driven*. Mais il nous semble que des renvois eussent été utiles entre, par exemple, *alpha-beta method* et *branch-and-bound algorithm*, puisque le deuxième terme désigne, en recherche opérationnelle, le même procédé que le premier.

La microstructure

La microstructure de MDAI est simple, conformément à une tradition illustrée particulièrement par les dictionnaires Elsevier : il s'agit d'une liste d'équivalents présentés en colonnes, une langue par ligne. Il y a peu de commentaires discursifs ou même d'ébauches de définitions. MDAI s'oppose en ceci à Pavel (1987), Otman (1991) et Kodratoff et Barès (1991), moins riches en effectif mais plus attrayants pour qui veut s'initier au domaine ou à un des sous-domaines.

Si l'on compare les équivalences qui sont proposées par MDAI avec celles de Pavel (1987) (P) et celle d'Otman (1991) (O) pour les mêmes entrées, on constate là encore un certain nombre de différences, qui reflètent bien le flou terminologique du domaine : *alpha-beta method* = *méthode alpha-béta* (O), *procédure alpha-béta* (P), *procédé alpha-béta* (MDAI) ; *automated vision* = *vision par ordinateur* (O), *vision artificielle* (P), *vision automatisée* (MDAI), etc.

Certaines équivalences françaises proposées par MDAI semblent discutables. Pour *explanation facility*, il renvoie à *explanation subsystem*, où l'on trouve *subystème d'explication* et *module d'explication*. Ce *sub-* est bien étrange : c'est la seule occurrence de ce préfixe, alors qu'il y a 21 occurrences de *sous-*, dont une d'ailleurs dans le terme *sous-système*. De même, *iterative deepening* est traduit par *profondissement itératif*, alors que

deepening (of a search) est traduit par *approfondissement*. Ce *profondissement* est très curieux, même s'il est... réitéré !

Le dictionnaire d'Otto Vollnhals rendra des services dans un domaine où la terminologie est bourgeonnante et en constante évolution. Il faut rendre hommage à l'auteur et à l'éditeur qui ont su combler une lacune. Ils n'ont pas cherché à attendre un début de stabilisation de la terminologie de l'IA pour procurer aux traducteurs, sous une forme agréable et facile à manier, l'information dont ils ont besoin *hic et nunc*. On peut bien sûr regretter que l'auteur n'ait pas proposé des explications de concepts qui semblent souvent bien obscurs pour le traducteur. Mais peut-être aurait-on alors changé d'outil : compte tenu du multilinguisme de cet ouvrage, sa taille l'aurait rendu difficile à utiliser.

Dans sa préface, Otman (1991) parle d'une version ASCII de son texte. Nous serions bien surpris si Vollnhals n'avait pas, lui aussi, en réserve, le même «prototypage» qu'un enrichissement définitionnel (ou doit-on simplement dire «notionnel» ?) rendrait très performant pour un traducteur. Après tout, il n'est pas anormal de souhaiter qu'un travail lexicographique sur l'intelligence artificielle soit utilisable sur des machines, ainsi rendues plus «intelligentes».

HENRI BÉJOINT ET PHILIPPE THOIRON
Université Lumière Lyon 2, Lyon, France

Note

1. On notera d'ailleurs que Kodratoff et Barès (1991) abordent la question de la manière suivante. Sous *Bayésien (modèle)*, défini comme «modèle de calcul des probabilités conditionnelles qui applique le théorème de Bayes», le théorème est énoncé et les auteurs donnent un exemple «concret» emprunté à la médecine. Il semble bien que, même dans le cadre de l'intelligence artificielle, c'est dans le domaine des mathématiques, et non dans le sous-domaine des probabilités, que doit être classé ce terme. Ceci ne préjuge en rien de l'utilisation qui en sera faite ; il est bien connu que les probabilités sont mises en œuvre dans bon nombre de théories.

RÉFÉRENCES

- KODRATOFF, Yves et Michel BARÈS (1991) : *Base terminologique de l'intelligence artificielle*, Paris, Technique et documentation, Lavoisier.
- OTMAN, Gabriel (1991) : *Vocabulaire de l'intelligence artificielle*, Paris, EC2.
- PAVEL, Silvia (1987) : *Vocabulaire de l'intelligence artificielle*, Ottawa, Secrétariat d'État.