

Université de Paris-Sud
Centre d'Orsay
DEUG MIAS
M3 — spécialité Informatique

Interprétation des langages de programmation

Jean-Pierre Jouannaud
Claude Marché
Ralf Treinen

Adresses des auteurs:

Jean-Pierre Jouannaud
L.R.I., Bâtiment 490
Université de Paris-Sud
91405 Orsay cedex
France

Tél : 01 69 15 69 05
Bureau 148
Email: Jean-Pierre.Jouannaud@lri.fr
WWW : <http://www.lri.fr/Francais/Recherche/demons/membres/jouannau.html>

Claude Marché
L.R.I., Bâtiment 490
Université de Paris-Sud
91405 Orsay cedex
France

Tél : 01 69 15 64 85
Bureau 118
Email: Claude.Marche@lri.fr
WWW : <http://www.lri.fr/~marche/>

Ralf Treinen
L.R.I., Bâtiment 490
Université de Paris-Sud
91405 Orsay cedex
France

Tél : 01 69 15 65 92
Bureau 152
Email: Ralf.Treinen@lri.fr
WWW : <http://www.lri.fr/~treinen/>

Table des matières

Introduction	1
I Syntaxe	5
1 Notion de langage	7
1.1 Ensemble des mots sur un alphabet donné	7
1.2 Langages	8
1.3 Notion intuitive de grammaire	8
1.4 Définition formelle des grammaires	9
1.5 Exercices	11
Méthode 1: par énumération des mots du langage.	15
Méthodes 2: analyse par cas en fonction de la première lettre du mot à reconnaître.	16
Méthode 2a	17
Méthode 2b	18
Complément de CAML	18
2 Grammaires régulières et automates	25
2.1 Grammaires régulières	25
2.2 Automates déterministes	27
2.3 Programmation des automates	28
2.3.1 Spécifications	28
2.3.2 Implantations	29
2.3.3 Fonctions <i>exécute</i> et <i>reconnaît</i>	29
2.3.4 Implantation du type automate	30
2.4 Automates non-déterministes et grammaires régulières	32
2.5 Déterminisation des automates	34
2.6 Minimisation des automates	35
2.7 Exercices	40
3 Expressions rationnelles et analyse lexicale	48
3.1 Automates non-déterministes avec transitions vides	49
3.2 Propriétés de clôture des langages réguliers	51
3.2.1 Concaténation et Puissance	53
3.2.2 Union	53
3.2.3 Itération et itération stricte	53
3.2.4 Complémentation et intersection	54
3.3 Expressions rationnelles	54
3.4 Analyse lexicale avec l'outil LEX	56
3.4.1 Notion de token	56

3.4.2	Fichier de description LEX	56
3.4.3	Description des tokens	57
3.4.4	Règles de priorité	57
3.5	Exercices	59
4	Analyse syntaxique	71
4.1	Dérivations gauches et droites	71
4.2	Arbres syntaxiques	73
4.3	Applications à l'analyse syntaxique	75
4.3.1	Analyse syntaxique descendante	75
4.3.2	Analyse syntaxique ascendante	77
4.4	Génération d'analyseurs syntaxiques avec YACC	79
4.5	Exercices	80
5	Syntaxe abstraite des langages	92
5.1	Grammaires abstraites	92
5.2	Arbres de syntaxe abstraite	93
5.3	Représentation des valeurs des tokens dans l'arbre	94
5.4	Calcul des arbres de syntaxe abstraite	96
5.5	Codages des arbres de syntaxe abstraite	96
5.6	Exercices	96
II	Sémantique	98
1	Généralités sur la sémantique des langages	100
1.1	Notion de jugement	100
1.2	Règles sémantiques	101
1.3	Notion de dérivation d'un jugement	103
1.4	Codage des règles sémantiques	104
1.5	Exercices	105
2	Sémantique des expressions arithmétiques et logiques	108
2.1	Expressions arithmétiques	108
2.1.1	Syntaxe abstraite	108
2.1.2	Évaluation	109
2.2	Expressions logiques	109
2.2.1	Syntaxe abstraite	109
2.2.2	Évaluation	110
2.3	Expressions arithmético-logiques	110
2.3.1	Syntaxe abstraite à deux sortes	111
2.3.2	Évaluation	111
2.3.3	Syntaxe abstraite à une sorte	113
2.3.4	Typage	113
2.3.5	Évaluation	115
2.3.6	Notion de typage fort	115
2.4	Expressions arithmétiques surchargées	116
2.4.1	Syntaxe abstraite	116
2.4.2	Typage	117
2.4.3	Évaluation	119
2.5	Exercices	119

3	Variables et environnements	120
3.1	Cas des expressions arithmétiques seules	120
3.1.1	Syntaxe abstraite	120
3.1.2	Typage	121
3.1.3	Évaluation	121
3.2	Les expressions arithmético-logiques	124
3.2.1	Typage	125
3.2.2	Évaluation	125
3.3	Exercices	128
4	Fonctions à un argument	131
4.1	Fonctions non récursives	131
4.1.1	Syntaxe abstraite	131
4.1.2	Typage	132
4.1.3	Évaluation	132
4.2	Fonctions récursives	133
4.2.1	Typage	134
4.2.2	Évaluation	134
4.3	Exercices	135
5	Sémantique d'un langage fonctionnel	139
5.1	Syntaxes concrète et abstraite	140
5.2	Typage	140
5.3	Évaluation	141
5.4	Fonctions à plusieurs arguments	141
5.5	Problème de la liaison dynamique	142
5.6	Codage des expressions <code>let</code>	142
6	Types structurés	144
6.1	Les types produits	144
6.1.1	Le type couple	144
	Syntaxe	144
	Typage	145
	Évaluation	145
6.1.2	Les types enregistrements	145
	Typage	146
6.2	Les types sommes	146
6.2.1	Somme anonyme de deux types	147
	Syntaxe	147
	Typage	147
	Évaluation	148
6.2.2	Sommes étiquetés	148
7	Sémantique d'un langage impératif	149
7.1	Syntaxe concrète et abstraite	149
7.2	Typage	150
7.3	Évaluation	151
7.4	Traduction des programmes impératifs en programmes fonctionnels récursifs	151
7.5	Exercices	155

III Études de cas	158
1 Généralités sur les études de cas	160
1.1 Le contenu des rapports	160
1.1.1 Manuel d'utilisation	160
1.1.2 Structure du programme et découpage en modules	160
1.1.3 Description des modules	160
1.1.4 Jeu de tests	161
1.1.5 Soutenance	161
2 Compilation d'un mini-Pascal	162
3 Un langage d'expressions rationnelles	163
4 Un langage pour la composition de documents	164
4.1 Documents simples	164
4.1.3 Le module POSTSCR, et l'outil GhostView	165
4.1 Ceci est un titre de section	166
4.2 Voici la deuxième section	166
4.2.1 Et ceci est une sous-section	166
4.2.2 Deuxième sous-section	166
4.1.4 Les fichiers sources et leur syntaxe	166
4.1.5 Extensions du langage	167
Numérotation des pages et entêtes	167
Symboles spéciaux	168
Paramétrage du style	168
4.1.6 Interface PASCAL du module POSTSCR	169
4.2 Les formules mathématiques	170
4.2.1 Syntaxe des formules	170
4.2.2 Indications pour la composition des formules	172
4.3 Les tableaux	173
4.3.1 Syntaxe	173
4.3.2 Indications pour la composition	174
4.4 Calculs automatiques	174
4.4.1 Syntaxe	174
4.4.2 Sémantique	175
5 Un navigateur hypertexte	176
5.1 Le langage HTML de base	176
5.2 Énumérations et tableaux	178
5.2.1 Les tableaux	179
5.2.2 Les énumérations	179
5.2.3 Remarques générales sur le langage	179
Bibliographie	183
Index	186

Introduction

La notion de langage est issue de la linguistique, qui est la science dont l'objet d'étude est le langage naturel, le français par exemple. Les linguistes cherchent à décrire les langues naturelles par des règles permettant d'engendrer les phrases correctes des langages en question. Ces mécanismes de description portent le nom familier de *grammaire*.

L'informatique a emprunté à la linguistique la notion de langage, ainsi que leur description par des grammaires. Les langages informatiques doivent toutefois satisfaire des critères contraignants auxquels les langages naturels ne sont pas assujettis, car les langages informatiques servent à décrire des calculs, et il est impératif qu'une phrase d'un langage informatique définisse un calcul bien précis, c'est-à-dire unique. Les langages naturels sont dits *ambigus*, à l'inverse des langages informatiques qui ne doivent pas l'être.

La notion de modèle de calcul est bien sûr au cœur de la discipline Informatique, même si elle fut à l'origine introduite par des mathématiciens dans les années 1930. Les modèles de calculs sont nombreux, il en existe au moins un par style de langage, impératif comme PASCAL, fonctionnel (on dit aussi applicatif) comme ML et ses dialectes, logique comme PROLOG, par contraintes comme CHIP, orienté objets comme JAVA. Ces modèles de calculs sont dits *universels* en ce sens qu'ils permettent de décrire n'importe quel calcul exécutable sur un ordinateur. La description et le calcul lui-même dépendent bien sûr du modèle utilisé.

Notre but dans ce cours n'est pas d'étudier les modèles de calcul précités, mais de comprendre quels sont les mécanismes permettant d'associer un calcul à un programme écrit dans un langage impératif comme PASCAL ou applicatif comme CAML. Il n'est pas non plus question d'aborder le mécanisme d'exécution d'un programme PASCAL ou CAML arbitraire, nous nous contenterons de fragments élémentaires de ces langages, et travaillerons à partir d'exemples simples, en particulier l'interprétation des expressions arithmétiques, qui a le mérite d'être proche du fonctionnement d'une calculette de poche de bas de gamme.

L'interprétation d'un programme se passe en plusieurs phases que nous considérerons comme étant bien distinctes de manière à en faciliter la description et la compréhension. Ces différentes phases sont visualisées sur la figure 1. Les deux premières phases sont dites *syntactiques*, elles ont pour but de vérifier la bonne utilisation des règles de construction des programmes PASCAL ou ML ; alors que les autres phases sont dites *sémantiques*, elles ont pour but d'associer à chaque programme le calcul qu'il définit.

La première phase est *l'analyse lexicale*. La donnée de l'analyse lexicale est une suite de caractères composant un programme. Cette suite de caractères se présente comme un flot continu sans structure, on parle de flot de caractères. Le rôle de l'analyse lexicale est de reconnaître des groupes de caractères appelés *unités lexicales* par analogie avec les langues naturelles où les mots forment des unités lexicales (lexique est d'ailleurs un synonyme de dictionnaire). Pour les langages informatiques, les unités lexicales sont les noms de variables, les constantes numériques comme 3.141592, les mots clés comme `while`, les opérateurs arithmétiques et logiques comme `+`, `or` ou `>=`. Le résultat de l'analyse lexicale est semblable à sa donnée, il s'agit d'un nouveau flot de caractères appelés *tokens*, dont les espacements ont en général été éliminés. L'analyse lexicale est en général faite à l'aide d'un outil originalement disponible sous les systèmes d'exploitation de la famille UNIX, le logiciel LEX. L'utilisation

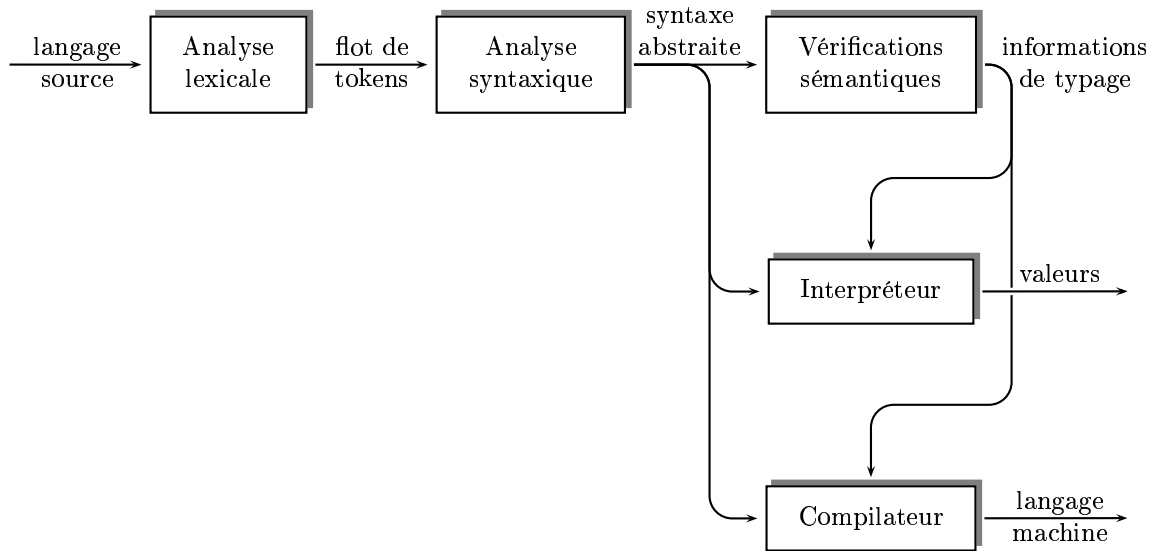


FIG. 1 – Les différentes phases de l'interprétation d'un langage

de LEX demande une connaissance superficielle des mécanismes mis en jeu, qui reposent sur les notions de grammaires régulières, automates et expressions rationnelles. Ces notions, et leur relations, seront abordées dans le cours par le biais de nombreux exemples simples.

La seconde phase est l'analyse syntaxique. La donnée est cette fois le flot de tokens résultat de l'analyse lexicale. Ce flot de tokens est transformé en un *arbre de syntaxe abstraite* qui décrit les différentes opérations à effectuer ainsi que leur enchaînement. L'analyse syntaxique est en général faite à l'aide d'un autre outil analogue à LEX, le logiciel YACC. L'utilisation de YACC demande une connaissance superficielle des mécanismes mis en jeu, qui reposent sur les notions de grammaire hors-contexte et d'analyse ascendante. Ces notions seront abordées dans le cours par le biais de l'exemple des expressions arithmétiques.

Les phases ultérieures ont pour but de *calculer* la signification d'un programme, on dit aussi sa *sémantique*, ce qui sera effectué à partir de l'arbre de syntaxe abstraite du programme, dont l'unicité sera donc une propriété essentielle. Ces calculs ont une structure très particulière, on parle d'homomorphisme, car ils respectent la structure de l'arbre de syntaxe abstraite. On dit aussi que la sémantique d'un programme est *compositionnelle*, car la sémantique du tout (l'arbre) est une fonction de la sémantique de ses parties (les sous-arbres). Le formalisme utilisé pour décrire ces calculs est celui des *règles sémantiques*, un formalisme très général dont le but est de décrire des calculs s'effectuant sur des arbres. Ce formalisme trouve ses racines en logique mathématique, ce qui explique en partie le vocabulaire que nous utiliserons.

La troisième phase (vérifications sémantiques) opère de nombreux contrôles, en particulier des contrôles de types et des déclarations des variables préalablement à leur utilisation, ainsi que des calculs simples. Ces contrôles sont décrits par des règles sémantiques particulières appelées *jugements de typage*. Les informations obtenues lors de ces calculs pourront servir à l'interprétation elle-même. Nous nous contenterons d'effectuer des contrôles de type, ce que nous décrirons une fois de plus par le biais d'exemples simples pris en particulier dans le domaine du typage des expressions arithmétiques.

La quatrième phase est celle de l'interprétation elle-même, dont le principe est d'exé-

cuter pas à pas les instructions du programme telles qu'elles sont représentées dans l'arbre syntaxique, à partir des données fournies par l'utilisateur. Cette phase est en fait la plus complexe, car elle implique la mise en œuvre de mécanismes de gestion de la mémoire. Ces mécanismes dépendent de la catégorie de variables présentes dans le programme (variables locales ou globales, variables simples ou structurées, variables de taille fixes ou non, tous renseignements qui se trouvent dans la table des symboles), de la possibilité pour le programme d'appeler d'autres programmes ou de s'appeler lui-même récursivement, etc. La gestion de la mémoire comporte plusieurs aspects : construction de l'environnement de liaison des variables (plusieurs copies d'une même variable peuvent coexister, correspondant à plusieurs sections ou copies d'un même programme) ; récupération des zones de mémoire inutilisées (ramassage de miettes), et compactage de la mémoire si nécessaire afin d'en éviter la fragmentation. Ces deux derniers aspects sont généralement omis dans un interprète, c'est-à-dire traités automatiquement par les mécanismes correspondants du langage d'implantation. Cette phase peut être faite par des outils dans les cas simples, mais nécessite en tout cas une bonne compréhension des mécanismes mis en œuvre. Il est hors de question de traiter en détail les mécanismes complexes de gestion de la mémoire nécessités par un langage même aussi simple que PASCAL ou CAML. Nous nous contenterons une fois de plus d'étudier l'interprétation de sous-ensembles simplifiés de ces langages, contenant en particulier le langage des expressions arithmétiques.

La compilation est un mécanisme encore plus complexe que l'interprétation, dont le but est de traduire le programme du *langage source* vers le langage de la *machine cible* de manière à obtenir des performances élevées. L'interprétation est utilisée dans au moins trois circonstances distinctes : pour l'exécution de langages prototypes, dont la définition est encore sujette à modification ; pour l'exécution de langages d'interfaces, pour lesquels la rapidité n'est pas un problème crucial ; lors de la mise au point des programmes, les interprètes se prêtant mieux que les compilateurs à la conception d'outils de déverminage sophistiqués.

Grammaires et automates ne sont pas des outils spécifiques de l'interprétation (ou la compilation) de programmes informatiques, même si leur étude a été en grande partie historiquement motivée par cette application. Les grammaires sont bien sûr au cœur de la linguistique, comme nous l'avons déjà dit. De leur côté, les automates ont envahi de nombreux champs de la science informatique, et même au delà. Ils sont utilisés dans des domaines aussi variés que la preuve de correction de circuits, la recherche de mots dans un dictionnaire, le traitement automatisé des langues naturelles, le traitement automatisé des images, la modélisation de certaines formes de jeux, le modèle d'exécution des langages parallèles synchrones, la théorie des groupes dits automatiques, la théorie des nombres, la théorie des réseaux cristallins, le traitement du génome, et bien sûr l'automatique. Enfin, le langage étant un moyen de communication privilégié, des langages ad'hoc sont également utilisés dans un très grand nombre de domaines non scientifiques. Un exemple intéressant est fourni par la lecture de ce polycopié : chaque définition est introduite par le mot clé **Définition** écrit en caractères gras, chaque figure est pourvue d'un titre, les énumérations sont numérotées de manière consistante, etc. Ce polycopié suit donc un certain nombre de règles typographiques que l'on retrouve à des variantes près, dans tous les textes mathématiques. De fait, ce polycopié a été écrit en LaTeX, un langage spécialisé dans l'écriture de documents. Le compilateur LaTeX prend un texte écrit par l'utilisateur et le compile en un texte Postscript. Postscript est un autre langage, en fait le langage machine des imprimantes modernes, en particulier les imprimantes laser. L'exécution du code Postscript engendré produit donc comme résultat l'impression du document.

Un dernier exemple est le langage HTML (Hyper Text Markup Language). Ce n'est pas un langage de programmation mais un langage de description d'hypertexte, c'est-à-dire de pages de texte accompagné de « liens » pointant sur d'autres pages de texte. C'est dans ce langage que sont écrites les pages du *World Wide Web* qui peuvent être ainsi visualisées sur

n'importe quelle type d'ordinateur.

Le document est structuré en deux parties : syntaxe, puis sémantique.

La première partie comporte cinq chapitres. Le chapitre 1 décrit les notions de mots et de langage formel sur un alphabet, ainsi que les opérations importantes sur les mots et les langages utilisées dans la suite. Il introduit les grammaires, un formalisme génératif pour décrire des langages. Le chapitre 2 étudie plus à fond les grammaires régulières et le mécanisme de reconnaissance des langages associés, les automates. Il débouche sur l'algorithme de minimisation des automates déterministes. Le chapitre 3 continue avec l'étude des expressions rationnelles, un outil commode de description des langages réguliers, et son application à l'analyse lexicale. Il se conclut avec une description rapide du logiciel LEX. Le chapitre 4 continue avec une étude plus approfondie des grammaires générales, en introduisant les notions de dérivations et d'arbre syntaxique, et débouche sur les principes généraux d'analyse syntaxique, et leur application dans le cadre du logiciel YACC. Enfin, le chapitre 5 décrit la notion de syntaxe abstraite, qui est à la base de la description de la sémantique des langages telle que nous l'exposons dans la seconde partie du document. Nous y étudions également comment calculer la syntaxe abstraite d'un langage à l'aide de YACC.

La seconde partie comporte cinq chapitres. Le chapitre 1, décrit la notion de règle sémantique, un formalisme adapté à la description de la sémantique des langages, de même que les grammaires sont adaptées à la description de leur syntaxe. Pour chaque langage traité, nous étudions successivement deux aspects : le typage, et l'évaluation, qui donnent lieu à ce que l'on appelle des jugements de typage et des jugements d'évaluation. Le chapitre 2 traite du langage des expressions arithmétiques et booléennes, et introduit les règles de typage. Le chapitre 3 traite les définitions de variables et donc les problèmes de portée, en introduisant la notion fondamentale *d'environnement*. Le chapitre 4 traite les expressions fonctionnelles, dans le cas particulier simple des fonctions à un unique argument, et aborde les problèmes liés à la récursivité. Enfin, le chapitre 5 décrit un langage fonctionnel simple qui réunit les divers aspects abordés séparément dans les chapitres précédents.

La bonne compréhension de ce cours nécessite une pratique rudimentaire des langages PASCAL et CAML ou des langages analogues, ainsi que la connaissance des notions de graphe et d'arbre étiqueté. Sinon, le cours ne demande aucune connaissance mathématique particulière. En particulier, et c'est délibéré, il ne contient aucune preuve des théorèmes ou propriétés énoncés. Le lecteur familier avec le raisonnement mathématique est toutefois invité à suppléer à ce manque.

Exercice 0.1

Trouver d'autres exemples d'utilisation de langages ad'hoc, dans les sciences traditionnelles (biologie, chimie, ...) ou dans des domaines non-scientifiques.

Correction : *Le code génétique en biologie, c'est bon ? La notation $C_5H_{17}O_{bidule}$ que je sais plus comment ça s'appelle en chimie, c'est bon aussi ? La langue de bois en politique ?*

Première partie

Syntaxe

Chapitre 1

Notion de langage

Le but premier de ce chapitre est définir la notion de langage. Cette notion est très simple, et commune à tous les champs disciplinaires qui l'utilisent. Un langage est tout simplement un ensemble fini ou infini de mots formés sur un certain alphabet, qui pour nous sera en général le jeu de caractères ASCII ou l'un de ses sous-ensembles. Le second but est d'introduire un procédé permettant de décrire des langages, les grammaires.

1.1 Ensemble des mots sur un alphabet donné

Définition 1.1 *Étant donné un ensemble fini V appelé vocabulaire ou alphabet dont les éléments sont appelés des symboles ou lettres, on appelle mot sur V toute suite finie $u = \{\alpha_i\}_{i \in [1..n]}$ d'éléments de V . n s'appelle la longueur du mot u et est notée $|u|$.*

D'une manière générale, dans ce cours, nous utiliserons des lettres grecques pour dénoter des symboles, et les lettres u , v et w pour dénoter des mots.

La notation d'un mot sous forme d'une suite n'étant pas très pratique, un mot $u = \{\alpha_i\}_{i \in [1..n]}$ sera en général noté simplement par juxtaposition de ses lettres : $u = \alpha_1 \cdots \alpha_n$, et le mot vide (de longueur nulle) sera alors noté par le symbole ε .

L'ensemble des mots sur V est muni d'une opération binaire de composition interne appelée *produit de concaténation*. Le produit de concaténation de deux mots u et v , noté $u \times v$ ou en abrégé uv (la notation $u \cdot v$ est aussi utilisée), est défini comme suit :

$$\{\alpha_i\}_{i \in [1..m]} \times \{\beta_i\}_{i \in [1..n]} = \{\gamma_i\}_{i \in [1..m+n]} \text{ tel que } \begin{cases} \gamma_i = \alpha_i, & \text{pour tout } i \in [1..m] \\ \gamma_{m+i} = \beta_i, & \text{pour tout } i \in [1..n] \end{cases}$$

Le produit de concaténation est bien sûr associatif : $(uv)w = u(vw)$, et possède un élément neutre, le mot vide : $u\varepsilon = \varepsilon u = u$.

Le produit de concaténation nous permet de définir par récurrence les *puissances* d'un mot u :

$$u^0 = \varepsilon, \quad u^{p+1} = u.u^p$$

La puissance p fois du mot $u = \{\alpha_i\}_{i \in [1..n]}$ est donc le mot $v = \{\beta_i\}_{i \in [1..np]}$ de longueur np tel que $\beta_{kn+i} = \alpha_i$, pour tout $k \in [0..p-1]$ et $i \in [1..n]$, ou encore, avec la notation par juxtaposition :

$$(\alpha_1 \cdots \alpha_n)^p = \underbrace{(\alpha_1 \cdots \alpha_n) \times \cdots \times (\alpha_1 \cdots \alpha_n)}_{p \text{ fois}}$$

1.2 Langages

Définition 1.2 On appelle langage sur un vocabulaire V tout ensemble de mots sur V .

Les opérations fondamentales sur les langages sont :

- le produit de concaténation : $L_1 \times L_2 = \{u \times v \mid u \in L_1, v \in L_2\}$, noté aussi $L_1 L_2$ en abrégé ;
- la puissance : $L^0 = \{\varepsilon\}$ et $L^{p+1} = L \times L^p$ c'est-à-dire $L^p = \{u_1 \cdots u_p \mid u_1, \dots, u_p \in L\}$;
- l'itéré : $L^* = \bigcup_{p \geq 0} L^p$ c'est-à-dire $L^* = \{u_1 \cdots u_p \mid p \geq 0, u_1, \dots, u_p \in L\}$;
- l'itéré strict : $L^+ = \bigcup_{p > 0} L^p$ c'est-à-dire $L^+ = \{u_1 \cdots u_p \mid p \geq 1, u_1, \dots, u_p \in L\}$.

L'opération d'itération est encore appelée *étoile de Kleene* du nom de son inventeur. Il est facile de voir que l'itéré de V est l'ensemble de tous les mots, ce dernier sera donc naturellement noté par V^* .

Les langages étant des ensembles, ils en héritent également les opérations booléennes habituelles :

- l'union : $L \cup L' = \{u \mid u \in L \text{ ou } u \in L'\}$;
- l'intersection : $L \cap L' = \{u \mid u \in L \text{ et } u \in L'\}$;
- le complémentaire : $\overline{L} = \{u \mid u \in V^* \text{ et } u \notin L\}$.

La question fondamentale qui va nous préoccuper dans ce cours, directement ou indirectement, est celle de la description d'un langage donné. La difficulté bien entendu, est qu'il faut pouvoir décrire des langages infinis, puisqu'il y a une infinité de programmes possible dans un langage de programmation. On distingue 3 grandes méthodes de description :

1. La méthode *généralisatrice* a pour principe la donnée d'un algorithme permettant d'engendrer l'ensemble des phrases du langage à décrire à partir de ce que l'on appelle une *grammaire*. Il s'agit donc en fait d'une description en extension adaptée au cas des langages infinis.
2. La seconde méthode consiste à donner un algorithme permettant de décider si une phrase donnée appartient ou pas au langage, à l'aide de ce que l'on appelle un *automate* (plus généralement une machine de Turing).
3. La dernière méthode, de nature algébrique, consiste à donner une notation spécifique, les *expressions rationnelles*, permettant de représenter un langage simple par une expression algébrique.

Nous aurons l'occasion de rencontrer ces trois types de méthodes. Dans la suite de ce chapitre, nous allons brièvement introduire la première.

1.3 Notion intuitive de grammaire

Les grammaires sont des outils pour décrire des langages, inspirés des grammaires utilisées pour les phrases du langage naturel, comme par exemple :

Une phrase simple est constituée d'un sujet suivi d'un verbe suivi d'une préposition suivie d'un complément. Le sujet est un prénom ou bien un pronom. Les prénoms sont à choisir parmi « Alain », « Béatrice » et « Charles ». Les pronoms sont à choisir parmi « Il » et « Elle ». Les prépositions sont à choisir parmi « à », « de », « au » et « en ». Les verbes sont à choisir parmi « va » et « vient ». Les compléments sont à choisir parmi « l'école » et « marché ».

« *Charles va à l'école* » est un exemple de phrase simple au sens ci-dessus. Cette description fait ressortir plusieurs phénomènes :

1. Le français est utilisé pour sa propre description, ce qui nécessite l'utilisation de guillemets afin de différencier le français en tant que langage de description du français en tant que langage décrit.
2. Le français n'est pas adapté en tant que langage de description, pas plus que toute autre langue naturelle, car trop imprécis. La première phrase de la description, en particulier, a plusieurs interprétations suivant que « suivi d'une préposition » se rapporte à « verbe » ou à « sujet suivi d'un verbe ».
3. les concepts utilisés sont en tout petit nombre : « est constitué », « suivi de », « est », « ou », « à choisir », et la possibilité de nommer.

Tout cela milite en faveur d'une notation spécialisée simple dotée d'une notion de variable. Cette notation consiste en fait à décrire un langage par des règles de grammaire. Pour notre exemple, on obtient :

<i>Phrase-simple</i>	→	<i>Sujet Verbe Préposition Complément</i>
<i>Sujet</i>	→	<i>Prénom Pronom</i>
<i>Prénom</i>	→	Alain Béatrice Charles
<i>Pronom</i>	→	Il Elle
<i>Verbe</i>	→	va vient
<i>Préposition</i>	→	à de au en
<i>Complément</i>	→	l'école marché

On peut tout aussi bien décrire des langages informatiques, puisque tel est l'un de nos buts. Voici par exemple une grammaire décrivant les expressions arithmétiques telles qu'on les trouve dans la plupart des langages de programmation (on se limite aux opérations d'addition et de multiplication) :

<i>Expression</i>	→	<i>Expression + Expression Expression × Expression Entier</i>
<i>Entier</i>	→	<i>Chiffre Chiffre-non-nul Chiffres</i>
<i>Chiffres</i>	→	<i>Chiffre Chiffre Chiffres</i>
<i>Chiffre</i>	→	0 <i>Chiffre-non-nul</i>
<i>Chiffre-non-nul</i>	→	1 2 3 4 5 6 7 8 9

On verra par la suite qu'une grammaire peut présenter des ambiguïtés qu'il faudra alors éliminer. Dans la suite de ce chapitre, nous nous contenterons en fait de donner la définition de la notion de grammaire, tout en rejetant leur étude et leur emploi aux chapitres ultérieurs.

1.4 Définition formelle des grammaires

Comme on vient de le voir, toute suite de lettres de l'alphabet latin n'est pas une phrase en français. Il en est de même des langages informatiques : une suite de lettres est un programme si elle respecte certaines règles bien précises. La description de ces règles se fait à l'aide d'une *grammaire*. De la discussion précédente, on peut dire qu'une grammaire est donnée par un ensemble de règles associant à un certain composant du langage un choix possible pour ses constituants.

Définition 1.3 Une grammaire G est un quadruplet (V_t, V_n, S, R) où

1. V_t est un ensemble fini de symboles dits terminaux, appelé vocabulaire terminal ;

$N \rightarrow 0$		$N \rightarrow 0$	
$N \rightarrow 1$		$N \rightarrow 1M$	
$N \rightarrow 1M$	$N \rightarrow 0 \mid 1 \mid 1M$	$N \rightarrow 0 \mid 1M$	
$M \rightarrow 0$	$M \rightarrow 0 \mid 1 \mid 0M \mid 1M$	$M \rightarrow \varepsilon$	$M \rightarrow 0 \mid 1M$
$M \rightarrow 1$		$M \rightarrow 0M$	$M \rightarrow \varepsilon \mid 0M \mid 1M$
$M \rightarrow 0M$		$M \rightarrow 1M$	
$M \rightarrow 1M$			

FIG. 1.1 – Deux grammaires décrivant les entiers écrits en binaire

2. V_n est un ensemble fini (disjoint de V_t) de symboles dits non-terminaux, appelé vocabulaire non terminal ;
3. $V = V_n \cup V_t$ est appelé le vocabulaire de la grammaire G ;
4. S est un non-terminal particulier appelé source, ou axiome, ou encore symbole d'entrée ;
5. R est un ensemble fini de règles de grammaires de la forme $N \rightarrow u$, où $N \in V_n$ et $u \in V^*$

Dans ce cours, on utilisera en général des mots en italique commençant par une majuscule pour les non-terminaux, et des lettres minuscules ou des symboles spéciaux (+, ×, etc.) pour les terminaux.

Remarque : la notion de grammaire formelle est en fait plus générale que celle que nous venons de définir qui suffit pour les besoins informatiques, mais pas par exemple pour les besoins de la linguistique : la définition ci-dessus définit en fait les grammaires « hors-contexte », mais comme nous n'en rencontrerons pas d'autres, nous omettons ce qualificatif.

À titre d'exemple, nous nous proposons de décrire l'ensemble des entiers naturels en représentation binaire, vus comme des suites non vides de 0 et de 1 ayant la propriété que seule la suite 0 peut commencer par un 0. Pour cela, nous allons utiliser une grammaire pour laquelle l'ensemble des symboles non-terminaux est $\{ N, M \}$, la source est le non-terminal N , l'ensemble des symboles terminaux est $\{0, 1\}$ et les règles sont décrites figure 1.1. Ce tableau présente en fait quatre ensembles de règles, dont le premier et le troisième sont conformes à la définition, alors que les règles de même membre gauche ont été regroupées dans le second et le quatrième en séparant les différents membres droits par le symbole « $\mid$ » signifiant « ou », conformément à une convention usuelle : $N \rightarrow u_1 \mid \dots \mid u_n$ abrège l'ensemble des règles $N \rightarrow u_1, \dots, N \rightarrow u_n$. Les quatre ensembles de règles de la figure 1.1 décrivent le même langage, celui des nombres entiers naturels écrits en binaire, mais le troisième est plus compact que le premier grâce à l'utilisation du symbole ε . Le premier et le second ensemble de règles doivent être considérés comme identiques, de même que le troisième et le quatrième, car ils ne diffèrent que par l'économie d'écriture juste mentionnée.

Une grammaire est un mécanisme permettant *d'engendrer* les phrases du langage, en réécrivant un mot $u \in V^*$ en un nouveau mot $v \in V^*$ par l'opération consistant à remplacer une occurrence (quelconque) d'un non-terminal N (quelconque) de V_n présent dans u par un membre droit de règle dont N est membre gauche. Ce processus est itéré à partir de la source jusqu'à l'élimination complète des non-terminaux. Par exemple, en utilisant la première grammaire définie dans la figure 1.1 :

$$N \rightarrow 1M \rightarrow 11M \rightarrow 110M \rightarrow 1101M \rightarrow 11010$$

et en utilisant la deuxième 1.1 :

$$N \rightarrow 1M \rightarrow 11M \rightarrow 110M \rightarrow 1101M \rightarrow 11010M \rightarrow 11010$$

Définition 1.4 *Étant donnée une grammaire $G = \{V_t, V_n, S, R\}$, on dit que le mot $u \in V^*$ se réécrit en le mot $v \in V^*$ dans G , et l'on note $u \rightarrow v$, si*

1. $u = w_1 N w_2$, où $w_1 \in V^*$, $N \in V_n$ et $w_2 \in V^*$;
2. $v = w_1 w w_2$, où $w \in V^*$;
3. $N \rightarrow w$ est une règle de R .

On dit que le mot $v \in V^*$ dérive du mot $u \in V^*$, dans la grammaire G , ce que l'on note par $u \rightarrow^* v$, s'il existe une suite finie $w_0, w_1, \dots, w_n$ de mots de V^* telle que $w_0 = u$, $w_i \rightarrow w_{i+1}$ pour tout $i \in [0, n-1]$, et $w_n = v$.

Nous avons vu ci-dessus que le mot $11010M$ dérive du mot $1M$ dans la troisième grammaire de la figure 1.1. Par convention, nous indiquerons le non-terminal réécrit dans un mot en le soulignant, comme dans

$$1\underline{M} \rightarrow 11\underline{M} \rightarrow 110\underline{M} \rightarrow 1101\underline{M} \rightarrow 11010M$$

Définition 1.5 *Le langage engendré par la grammaire G est l'ensemble des mots de V_t^* qui dérivent de l'axiome de G , que l'on note par $\text{Lang}(G)$.*

Il est aisé de se convaincre que l'ensemble des mots qui appartiennent aux langages engendrés par les grammaires de la figure 1.1 est dans les deux cas l'ensemble des entiers naturels écrits en binaire, c'est-à-dire les suites de 0 et de 1 ne commençant pas par 0, sauf la suite 0 elle-même.

1.5 Exercices

Exercice 1.1

Soit l'alphabet $V = \{a, b\}$ et les mots $u = ab$ et $v = aba$. Calculer uv , vu , u^2 , u^3 et v^2 .

Correction : $uv = ababa$, $vu = abaab$, $u^2 = abab$, $u^3 = ababab$, $v^2 = abaaba$.

Exercice 1.2

Montrer que pour tout mot $u \in V^*$ et tous entiers $p, q \in \mathbb{N}$ on a $u^p u^q = u^{p+q} = u^q u^p$.

Exercice 1.3

Considérons les deux langages sur l'alphabet $\{a, b\}$ définis par $L_1 = \{a, ba\}$ et $L_2 = \{\varepsilon, b, aa\}$. Calculer $L_1 L_2$, $L_2 L_1$, L_1^2 , L_2^2 , L_1^3 , L_2^3 , L_1^* , L_2^* .

Correction :

$$\begin{aligned} L_1 L_2 &= \{a, ba, ab, bab, aaa, baaa\} \\ L_2 L_1 &= \{a, ba, ba, bba, aaa, aaba\} \\ L_1^2 &= \{aa, aba, baa, baba\} \\ L_2^2 &= \{\varepsilon, b, aa, bb, baa, aab, aaaa\} \\ L_1^3 &= \{aaa, abaa, baaa, babaa, aaba, ababa, baaba, bababa\} \\ L_2^3 &= \{\varepsilon, b, aa, bb, baa, aab, aaaa, bbb, baab, aabb, aaaab, bbaa, baaaa, aabaa, aaaaaa\} \end{aligned}$$

L_1^* est l'ensemble des mots sur $\{a, b\}$ tels qu'un b est toujours suivi d'un a , L_2^* est l'ensemble des mots sur $\{a, b\}$ tels que les a apparaissent toujours par deux.

Exercice 1.4

Correction : Le but du premier exercice est de familiariser les étudiants avec la manipulation des langages. Les questions 1 et 2 de cet exercice sont étroitement imbriquées et, on peut les traiter en même temps pour chacun des langages.

On considère les langages définis sur le vocabulaire $V = \{a, b\}$ par $L_1 = \{a\}^* \cdot \{b\}^*$, $L_2 = \{\{a\} \cup \{b\}\}^*$, $L_3 = \{ab\}^*$ et $L_4 = \{a\}^* \cup \{b\}^*$.

1. Trouver la forme générale des mots de chacun de ces langages.

Correction :

- L_1 est l'ensemble des mots de la forme $a^n b^m$, $n, m \geq 0$.
- L_2 est l'ensemble des mots formés à l'aide des lettres de V soit V^* .
- L_3 est l'ensemble des mots de la forme $\{ab\}^n$, $n \geq 0$.
- L_4 est l'union des mots de la forme a^n , $n \geq 0$ et des mots de la forme b^n , $n \geq 0$.

2. Les mots suivants appartiennent-ils à chacun de ces langages: a , b , aa , ab , ba , $abab$.

Correction : La réponse est évidente pour le langage L_2 qui contient tous les mots formés à l'aide des lettres du vocabulaire V .

Le tableau suivant résume les dérivations qui permettent de décider lorsque le mot appartient au langage

	a	b	aa	ab	ba	$abab$
L_1	$a^1 b^0$	$a^0 b^1$	$a^2 b^0$	$a^1 b^1$	Non	Non
L_3	Non	Non	Non	$\{ab\}^1$	Non	$\{ab\}^2$
L_4	a^1	b^1	a^2	Non	Non	Non

Les résultats négatifs sont les plus difficiles justifier. La non appartenance de ba et $abab$ à L_1 est due au fait que dans un mot du langage L_1 un a ne peut apparaître après un b . Les mots de L_3 sont une alternance stricte de lettre a et de lettre b commençant par un a . Les mots a , b , aa et ba ne répondent pas à cette définition. Enfin Les mots de L_4 ne contiennent que des a ou que des b . Les mots qui contiennent des a et des b n'appartiennent donc pas à L_4 .

3. Comparez ces langages (égalités, inclusions).

Correction :

- L_2 et L^* . Le langage L_2 est l'ensemble de tous les mots sur l'alphabet $\{a, b\}$. Il contient donc les langages L_1 , L_3 et L_4 . Nous avons donc $L_2 \supseteq L_1$, $L_2 \supseteq L_3$, $L_2 \supseteq L_4$.
- L_1 et L_3 . Nous avons qu'il existe au moins un mot (aa) qui appartient à L_1 et pas à L_3 et un mot ($abab$) qui appartient à L_3 et pas à L_1 . La comparaison entre L_3 et L_4 est elle aussi vite faite à l'aide des mots a et ab . En fait les seuls mots commun à ces deux langages sont ε et ab .
- L_1 et L_4 . Nous avons le mot ab qui appartient à L_1 mais pas à L_4 . Par contre tous les mots de L_4 sont aussi des mots de L_1 : Pour un mot de L_4 il y a deux possibilités. Soit ce mot est de la forme a^n avec $n \geq 0$ et dans ce cas il est aussi dans le langage L_1 sous la forme $a^n \cdot b^0$. Soit ce mot est de la forme b^n avec $n \geq 0$ et dans ce cas il est aussi dans le langage L_1 sous la forme $a^0 \cdot b^n$. Nous avons donc $L_4 \subseteq L_1$.

Exercice 1.5

Montrer que pour tous langages L , L_1 , L_2 et L_3 sur un vocabulaire V on a

1. $(L_1 \cup L_2)L_3 = (L_1 L_3) \cup (L_2 L_3)$;
2. $L_1(L_2 \cup L_3) = (L_1 L_2) \cup (L_1 L_3)$;
3. $L^* = L^+ \cup \{\varepsilon\}$.

4. $(L^*)^* = L^*$.

Exercice 1.6

On considère le langage L_0 sur l'alphabet $V_t = \{a, b, c, d\}$ défini par la grammaire suivante :

$$L \rightarrow a \mid b \mid cL \mid dLL$$

1. Montrer que ca , dba , $cdccab$ sont dans le langage L_0 .

Correction : Pour cette question, il suffit d'appliquer les règles de définition. Par la règle (1), $a \in L_0$, donc par (3), $ca \in L_0$. Par (2), $b \in L_0$ et $a \in L_0$ donc par (4), $dba \in L_0$. On a vu que $ca \in L_0$, donc par (3), $cca \in L_0$, et par (2) $b \in L_0$ donc par (4), $dccab \in L_0$, donc par (3) $cdccab \in L_0$.

2. Trouver tous les mots de L_0 d'au plus quatre lettres.

Correction : 1 lettre : a, b . 2 lettres : ca, cb . 3 lettres : $cca, ccb, daa, dab, dba, dbb$. 4 lettres : $ccca, cccb, cdaa, cdab, cdba, cdbb, daca, dacb, dbca, dbcb, dcaa, dcab, dcba, dcbb$.

3. Donner une méthode pour déterminer l'ensemble E_n des mots de L_0 ayant n lettres, à partir des mots de moins de n lettres. En déduire une formule de récurrence entre les E_n .

Correction : La méthode générale pour déterminer tous les mots de L_0 de n lettres est la suivante: Si $n = 1$ alors on a a et b . Si $n > 1$ alors de prendre tous les mots de $n - 1$ lettres et d'ajouter un c devant, et pour chaque i entre 1 et $n - 2$, de chercher tous les mots w_1 à i lettres et les mots w_2 à $n - i - 1$ lettres et de construire dw_1w_2 .

Nous avons donc la formule suivante mettant en relation les langages $E_i = \{u \in L_0 \mid |u| = i\}$:

$$\begin{aligned} E_1 &= \{a, b\} \\ E_{n+1} &= cE_n \cup \bigcup_{i=1}^{n-2} dE_i E_{n-i-1} \end{aligned}$$

Ces formules donnent immédiatement une fonction récursive qui étant donné un entier n calcule l'ensemble E_n .

4. En déduire une méthode pour décider si un mot est dans L_0 ou non. Que peut-on dire de la complexité des calculs effectués par cette méthode?

Correction : Une méthode pour décider si un mot u appartient à L_0 peut consister en tester l'appartenance de u à $E_{|u|}$. Ceci donne un programme très complexe en occupation mémoire et par conséquent en temps également.

5. Montrer que $adccba$ n'est pas dans le langage L_0 . On pourra raisonner sur la première lettre de ce mot.

Correction : Montrer qu'un mot n'appartient pas au langage est bien sûr beaucoup plus difficile. Pour montrer que ce mot n'est pas dans le langage, nous proposons de montrer que pour tout mot $w \in L_0$ de longueur au moins 2, la première lettre de w est un c ou un d .

En effet, si $w \in L_0$ est de longueur au moins 2, par la forme de la dernière production de la grammaire il doit exister soit un mot $w' \in L_0$ tel que $w = cw'$, soit deux mots w_1 et w_2 tels que $w = dw_1w_2$. Dans les deux cas la première lettre de w est c ou d .

6. Montrer que $cdbadc$ n'est pas dans le langage L_0 . On pourra raisonner sur la dernière lettre de ce mot.

Correction : Nous proposons de montrer que pour tout mot $w \in L_0$, la dernière lettre de w est un a ou b . On procède par récurrence sur la longueur de w . (Une alternative est de procéder par récurrence sur la construction du mot.)

cas $n = 1$: les mots de longueur 1 de L_0 sont a et b et vérifient la propriété voulue.

cas général : si $w \in L_0$ est de longueur au moins 2, par la forme de la dernière production de la grammaire il doit exister soit un mot $w' \in L_0$ tel que $w = cw'$, soit deux mots w_1 et w_2 tels que $w = dw_1w_2$. Par récurrence sur w' ou w_2 , la dernière lettre est un a ou un b .

7. Montrer la propriété suivante : si $w_1v_1 = w_2v_2$ avec w_1, w_2 dans L_0 , alors $w_1 = w_2$ et $v_1 = v_2$. En déduire que, si $w = w_1v_1$ avec w et w_1 dans L_0 , alors $v_1 = \varepsilon$.

Correction : Pour la deuxième propriété on prend $w_2 = w$ et $v_2 = \varepsilon$.

On montre la première propriété par récurrence sur la longueur n de w_1 :

cas $n = 0$: n'est pas possible parce que $\varepsilon \notin L_0$.

cas $n = 1$: alors $w_1 = a$ ou b , et nécessairement $w_1 = a$ resp. $w_2 = b$.

cas $n > 1$ et w_1 commence avec c : On a $w_1 = cw'_1$ avec $w'_1 \in L_0$, et également $w_2 = cw'_2$ avec $w'_2 \in L_0$. Donc, $cw'_1v_1 = cw'_2v_2$, alors $w'_1v_1 = w'_2v_2$. Par récurrence, $w'_1 = w'_2$, donc $w_1 = w_2$.

cas $n > 1$ et w_1 commence avec d : On a $w_1 = dw'_1w''_1$ avec $w'_1, w''_1 \in L_0$, et également $w_2 = dw'_2w''_2$ avec $w'_2, w''_2 \in L_0$. Donc, $dw'_1w''_1v_1 = dw'_2w''_2v_2$, alors $w'_1w''_1v_1 = w'_2w''_2v_2$. Par récurrence, $w'_1 = w'_2$, donc $w''_1v_1 = w''_2v_2$. On applique la récurrence à la dernière équation et obtient $w''_1 = w''_2$. En résumé, on a $w'_1w''_1 = w'_2w''_2$, c'est-à-dire $w_1 = w_2$.

On peut remarquer l'utilisation de l'associativité de la concaténation, et qu'il faut appliquer la récurrence deux fois.

8. Soit x_n le nombre de mots à n lettres de L_0 . Déterminer une relation de récurrence satisfaite par la suite $(x_n)_{n \in \mathbb{N}}$.

Correction : On remarque que grâce à la propriété précédente, la méthode de la question 3 ne peut pas engendrer un mot de deux manières différentes : si $w_1w_2 \in L_0$ et $w'_1w'_2 \in L_0$ alors $w_1 = w'_1$ et $w_2 = w'_2$.

Ainsi, on a $x_0 = 0$, $x_1 = 2$ et

$$x_n = x_{n-1} + \sum_{i=1}^{n-2} x_i \times x_{n-i-1}$$

pour $n \geq 2$.

9. Montrer que si w est dans L_0 , alors $|w|_a + |w|_b - |w|_d = 1$. Trouver un mot w , commençant par c ou d et finissant par a ou b , ne vérifiant pas cette propriété. Ce mot est-il dans L_0 ? Existe-il un mot w commençant par c ou d et finissant par a ou b , vérifiant la propriété précédente, mais n'étant pas dans L_0 ?

Correction : Encore par récurrence:

cas $n = 1$: $|a|_a + |a|_b - |a|_d = |b|_a + |b|_b - |b|_d = 1$.

cas général : si $w \in L_0$ est de longueur au moins 2, par la forme de la dernière production de la grammaire il doit exister soit un mot $w' \in L_0$ tel que $w = cw'$, soit deux mots w_1 et w_2 tels que $w = dw_1w_2$. Dans le premier cas, $|w|_a + |w|_b - |w|_d = |cw'|_a + |cw'|_b - |cw'|_d = |w'|_a + |w'|_b - |w'|_d = 1$ par récurrence. Dans le second cas, $|w|_a + |w|_b - |w|_d = |dw_1w_2|_a + |dw_1w_2|_b - |dw_1w_2|_d = |w_1|_a + |w_2|_a + |w_1|_b + |w_2|_b - 1 - |w_1|_d - |w_2|_d = 1 + 1 - 1 = 1$.

$w = caa$ commence bien par c ou d et finit bien par a ou b , et $|w|_a + |w|_b - |w|_d = 2$, il répond donc à la première question. Il n'est bien sûr pas dans L_0 puisqu'il ne vérifie la propriété.

$cadb$ répond à la deuxième question (on a vu à la question 2 qu'il n'était pas dans L_0).

Cette question et les précédentes doivent montrer aux étudiants la difficulté de trouver une méthode pour décider si un mot est dans le langage ou pas.

Correction : Attention pour le groupe S3, ceux-ci ne font pas l'étude de cas en parallèle et donc ne bénéficient pas des compléments de CAML. C'est pourquoi quelques compléments sont présentés au cours du TP. Avantage néanmoins pour les S3 : ils ont fait du CAML au semestre précédent donc on peut supposer qu'ils s'en souviennent mieux. Attention, dans le cours de CAML de M1, les étudiants ont appris à écrire les fonctions uniquement sous la forme

```
let f = function(x1,...,xn) -> ...
```

Il faudra les habituer très progressivement à écrire

```
let f = function x1 .. xn -> ...
```

ou même

```
let f x1 .. xn =
```

Exercice 1.7

D'une manière générale, un analyseur syntaxique pour un langage L est une fonction dont la spécification est la suivante :

analyse_L : chaîne de caractères $\rightarrow$ booléen

{ $\text{analyse}_L(w)$ retourne vrai si le mot w appartient au langage L , faux sinon }

On reprend l'exercice 1.6. Nous allons écrire une fonction analyse_{L_0} par trois méthodes différentes.

Méthode 1 : par énumération des mots du langage.

1. Écrire une fonction `prefixe` qui, étant donnés un mot w et une liste de mots $[w_1; \dots; w_n]$, retourne la liste $[ww_1; \dots; ww_n]$. Rappels de CAML : `[]` représente la liste vide, `::` l'opération d'ajout d'un élément en tête d'une liste, `^` l'opération de concaténation des chaînes de caractères. On définira `prefixe` par pattern-matching sur la liste en argument.

Correction :

```
let rec prefixe = function
  (w, []) -> []
| (w, x::r) -> (w^x)::(prefixe r)
;;
```

2. Écrire une fonction `prefixe_c` qui, étant donnée une liste de mots $[w_1; \dots; w_n]$, retourne la liste $[cw_1; \dots; cw_n]$. On réutilisera la fonction `prefixe`.

Correction :

```
let prefixe_c = function(l) => prefixe("c",l);;
```

3. Écrire une fonction `prefixe_d` qui, étant donnée deux listes $[v_1; \dots; v_n]$ et $[w_1; \dots; w_m]$, retourne la liste de tous les mots dv_iw_j avec $1 \leq i \leq n$ et $1 \leq j \leq m$. On définira cette fonction par récurrence sur la première liste. Rappel de CAML : `@` est la concaténation des listes. On réutilisera la fonction `prefixe`.

Correction :

```
let rec prefixe_d = function
  ([], l2) -> []
| (x::l1, l2) -> prefixe("d"^x, l2) @ prefixe_d(l1, l2);;
```

4. En utilisant la formule de récurrence trouvée à la question 3 de l'exercice 1.6, implanter une fonction `engendre` qui, étant donné un entier n , retourne l'ensemble E_n de tous les mots de L_0 de longueur n . On définira une fonction auxiliaire `union_partielle` locale à `engendre`, qui calcule

$$\text{union_partielle}(k) = \bigcup_{i=1}^k d.E_i.E_{n-i-1}$$

par récurrence sur k . On réutilisera également les fonctions `prefixe_c` et `prefixe_d`.

Correction :

```
let rec engendre = function
  0 -> []
| 1 -> ["a"; "b"]
| n ->
  let union_partielle = function(k) ->
    if k = 0
    then []
    else prefixe_d(engendre k, engendre(n-k-1)) @ union_partielle(k-1)
  in
  prefixe_c(engendre(n-1)) @ union_partielle(n-2);;
```

5. Écrire une fonction `analyse_L0` utilisant `engendre`. Complément de CAML : on utilisera les fonctions CAML `string_length` qui donne la longueur d'une chaîne de caractères, et `mem` telle que `(mem x l)` retourne `true` si l'élément x appartient à la liste l . Attention, pour utiliser cette fonction on écrira bien `(mem x l)` et pas `mem(x, l)`. Cette façon d'écrire l'appel à une fonction est commun à toutes les fonctions prédéfinies de CAML, et la différence entre les deux façons d'écrire sera expliquée plus tard en cours.

Correction :

```
let analyse_L0 = function(w) -> (mem w (engendre (string_length w)));;
```

Remarque : ne pas oublier les parenthèses autour de l'appel à engendre.

6. Tester `analyse_L0` sur des exemples de plus en plus long. Que peut-on dire sur l'efficacité de cette méthode? Comment augmente (approximativement) le temps de calcul en fonction de la longueur du mot à reconnaître?

Correction : Impraticable, si la longueur du mot est ≥ 7 . La complexité est exponentielle (au moins!).

Méthodes 2 : analyse par cas en fonction de la première lettre du mot à reconnaître. L'idée générale des deux méthodes que nous allons mettre en œuvre maintenant est la suivante : soit w un mot sur l'alphabet $\{a, b, c, d\}$. Si $w = \varepsilon$ alors $w \notin L_0$, sinon $w = \alpha w'$, et suivant le cas sur α :

- si $\alpha = a$ ou $\alpha = b$, alors $w \in L_0$ si et seulement si $w' = \varepsilon$.
- si $\alpha = b$, alors $w \in L_0$ si et seulement si $w' \in L_0$.
- si $\alpha = d$, alors $w \in L_0$ si et seulement si il existe deux mots w_1 et w_2 tels que $w' = w_1 w_2$ avec $w_1 \in L_0$ et $w_2 \in L_0$.

Méthode 2a *Correction* : En cas de manque de temps, on pourra passer directement à la méthode 2b

1. Écrire deux fonctions prenant en argument une chaîne de caractères : `premier(w)` retourne la première lettre de w , et `reste(w)` retourne le reste de w . On utilisera la fonction CAML `sub_string` telle que `(sub_string w n m)` retourne la sous-chaîne de w commençant au n -ième caractère et de longueur m . Attention, les caractères sont numérotés à partir de 0.

Correction :

```
let premier = function(s) -> sub_string s 0 1;;
let reste = function(s) -> sub_string s 1 ((string_length s)-1) ;;
```

2. Implanter la fonction `analyse_L0` par la méthode proposée. On écrira une fonction auxiliaire `cherche_un_decoupage` locale à `analyse_L0`, telle que `cherche_un_decoupage(w,n)` retourne `true` s'il existe deux mots w_1 et w_2 tels que $w = w_1.w_2$, $w_1 \in L_0$, $w_2 \in L_0$ et $|w_1| \leq n$. On raisonnera par récurrence sur n .

Correction :

```
let rec analyse_L0 = function
  "" -> false
| s ->
  let c = premier(s)
  and r = reste(s)
  in
    match c with
    "a" -> r=""
  | "b" -> r=""
  | "c" -> analyse_L0(r)
  | "d" ->
      let rec cherche_un_decoupage = function(w,n) ->
        if n=0
        then false
        else
          if analyse_L0(sub_string w 0 n) &
             analyse_L0(sub_string w n ((string_length w)-n))
          then true
          else cherche_un_decoupage(w,n-1)
      in cherche_un_decoupage(r,(string_length r))
  | _ -> false
;;
```

en PASCAL:

3. Que peut-on dire sur l'efficacité de cette deuxième méthode? Tester sur le mot $w_0 = dddadaaddaadaa$, puis sur dw_0w_0 et enfin $ddw_0w_0dw_0w_0$.

Correction : Beaucoup mieux mais encore exponentiel en fonction de la longueur du mot.

Méthode 2b On cherche à éviter la recherche d'un découpage de w en w_1w_2 dans le dernier cas, en exploitant la propriété montrée à la question 7 de l'exercice I.

1. Montrer la propriété suivante : pour tout mot w , il existe au plus un préfixe w_1 (c'est-à-dire $w = w_1w_2$) tel que $w_1 \in L_0$.
2. Réaliser une fonction `recherche_prefixe`, qui recherche un tel préfixe, dont la spécification exacte est

`recherche_prefixe`: chaîne de caractères $\rightarrow$ chaîne de caractères
 { `recherche_prefixe(w)` retourne l'unique mot w' tel que $w = w_1w'$ où $w_1 \in L_0$,
 déclenche l'erreur `failwith("pas de prefixe")` si aucun préfixe n'est dans L_0 }

Correction :

```
let rec recherche_prefixe = function
  "" -> failwith("pas de prefixe")
| s ->
  let c = premier(s)
  and r = reste(s)
  in
    match c with
    | "a" -> r
    | "b" -> r
    | "c" -> recherche_prefixe(r)
    | "d" -> recherche_prefixe(recherche_prefixe(r))
    | _ -> failwith("pas sur le bon alphabet")
;;
```

3. En déduire une nouvelle réalisation de `analyse_L0`.

Correction :

```
let analyse_L0(w) = recherche_prefixe(w)="";;
```

Cette fonction lève une exception quand le mot n'est pas dans le langage, on verra deux questions plus loin comment arranger cela.

4. Que peut-on dire sur l'efficacité de cette dernière méthode? Tester à nouveau sur les mots exemples de la méthode 2a.

Correction : *Encore mieux. C'est cette fois linéaire.*

Complément de CAML La fonction `analyse_L0` réalisée ici déclenche une erreur au lieu de retourner `false`. On dit qu'une *exception a été levée*. Nous allons maintenant apprendre comment *traiter* ces cas d'exception. Pour cela, CAML attend que le programmeur déclare des noms d'exception pour qualifier les cas d'erreur qui l'intéresse. Ceci s'effectue par l'instruction

```
exception nom d'exception;;
```


Ensuite, le programmeur peut lever ces exceptions par la fonction `raise` appliquée au nom de l'exception.

Par exemple, voici une fonction `deuxieme` qui retourne le deuxième élément d'une liste, ou bien lève l'exception `Liste_trop_courte` si la liste est trop courte :

```
exception Liste_trop_courte;;
let deuxieme = function
  x::(y::l) -> y
| _ -> raise(Liste_trop_courte);;
```

5. Refaire la réalisation de la fonction `recherche_prefixe` en levant l'exception `Erreur_lexicale` quand une lettre autre que *a*, *b*, *c* ou *d* est rencontrée; et `Erreur_syntaxique` quand aucun préfixe dans L_0 n'a pu être trouvé.

Correction :

```
exception Erreur_lexicale;;
exception Erreur_syntaxique;;

let rec recherche_prefixe = function
  "" -> raise(Erreur_syntaxique)
| s ->
  let c = premier(s)
  and r = reste(s)
  in
  match c with
  "a" -> r
| "b" -> r
| "t" -> recherche_prefixe(r)
| "c" -> recherche_prefixe(recherche_prefixe(r))
| _ -> raise(Erreur_lexicale)
;;
```

Lorsqu'une fonction f peut lever une certaine exception e , il est possible de *capturer* cette exception à l'aide de la construction suivante :

```
try expression contenant un appel à f
with e -> valeur de remplacement
```

Cette construction permet de dire que lorsque l'appel à la fonction f ne retourne pas un résultat normal mais l'exception e , alors la valeur de remplacement est à prendre à la place de l'expression contenant l'appel à f . Par exemple, voici une fonction qui utilise `deuxieme`, qui étant donnée une liste d'entiers donne la valeur de son deuxième élément plus 1, ou bien 0 si cette liste est trop courte :

```
#let f = function(l)->
  try deuxieme(l)+1
  with Liste_trop_courte -> 0;;
```

L'utilisation de `f` sous CAML donnera alors :

```
#f([1;2]);;
- : int = 3
#f([1]);;
- : int = 0
```

6. Sur ce principe, refaire la réalisation de `analyse_L0` en capturant l'exception `Erreur_syntaxique`. Ainsi, cette fonction retournera `true` pour un mot de L_0 , `false` pour un mot sur le bon alphabet mais pas dans L_0

Correction :

```
let analyse_L0 = function(w) ->
  try recherche_prefixe(w)=""
  with Erreur_syntaxique -> false
;;
```

Exercice 1.8

Décrire les langages engendrés par les grammaires sur $V_t = \{0, 1\}$, $V_n = \{N, M\}$, de source N et dont les règles sont les suivantes.

1. $N \rightarrow 0 \mid 1M$ $M \rightarrow 0 \mid 0M \mid 1M$

Correction : C'est l'ensemble des mots non vides se terminant par 0, et commençant par 1 excepté 0 lui-même. Il s'agit donc de l'ensemble des entiers pairs écrits en binaire.

2. $N \rightarrow 1M$ $M \rightarrow \varepsilon \mid 0N \mid 1M$

Correction : Cette fois c'est les impairs.

Exercice 1.9

Donner une grammaire décrivant les langages suivants :

- les entiers relatifs, le signe étant optionnel pour les entiers naturels. Cette grammaire devra par exemple reconnaître 0, 33, +20 et -13.
- les nombres réels en virgule flottante. Cette grammaire devra par exemple reconnaître 0.0, -1.33 et +20.13.
- les identificateurs des langages de programmation classiques, formés par des mots non vides sur l'alphabet des lettres majuscules ou minuscules, des chiffres et du caractère souligné `_` ; et commençant par une lettre. Cette grammaire devra par exemple reconnaître `x`, `Y1`, `ptr_v` et `Date`.
- les nombres réels en virgule flottante en acceptant les exposants. Cette grammaire devra reconnaître les réels de la question 2 ainsi que, par exemple, 1.0e3, -1.33E-3 et +20.13e+4.

Exercice 1.10

Donner des grammaires pour les langages suivants :

- Les mots de longueur paire sur l'alphabet $\{a, b\}$.
- Les mots sur l'alphabet $\{0, 1, 2, 3, 4, 5, 6, 7, 8, 9, -, :, \}$ affichés sur l'écran d'une montre digitale, dans le format année-mois-jour-heures:minutes:secondes.
- Les suites de a et de b ne contenant pas deux a consécutifs.
- Les suites de a et de b ne contenant pas deux a consécutifs ni deux b consécutifs.
- $\{a^n b^n \mid n \in \mathbb{N}\}$, où a^n indique la répétition n fois de la lettre a .
- $\{a^n b^m \mid n, m \in \mathbb{N}, n \leq m\}$

- l'ensemble des *palindromes* sur $\{a,b\}$, c'est-à-dire les mots qui se lisent indifféremment à l'envers ou à l'endroit (comme « radar » en français).

Correction :

1. $S \rightarrow \varepsilon \mid aaS \mid abS \mid baS \mid bbS$
- 2.
3. $S \rightarrow \varepsilon \mid bS \mid aB$, $B \rightarrow \varepsilon \mid bS$
4. $S \rightarrow \varepsilon \mid A \mid B$, $A \rightarrow \varepsilon \mid aB$, $B \rightarrow \varepsilon \mid bA$
5. $S \rightarrow \varepsilon \mid aSb$
6. $S \rightarrow aSb \mid T$, $T \rightarrow \varepsilon \mid bT$
7. $S \rightarrow \varepsilon \mid a \mid b \mid aSa \mid bSb$

Exercice 1.11

On note L le langage engendré par la grammaire sur l'alphabet $\{0,1\}$, d'axiome N , définie par les règles

$$N \rightarrow 11 \mid 1001 \mid N0 \mid NN$$

1. Trouver tous les mots de L de longueur inférieure ou égale à 5.

Correction : 11, 110, 1001, 1100, 1111, 10010, 11000, 11110, 11011

2. Si w est un mot sur $\{0,1\}$, on note $b(w)$ l'entier dont w est l'écriture binaire. Montrer que pour tout mot w de L , $b(w)$ est divisible par 3.

Correction : Par récurrence sur la longueur du mot. Si $|w| = 2$: OK. Si $|w| = n > 2$ et si l'on suppose que tout mot de L de longueur inférieure à n a un $b()$ divisible par 3, alors suivant la règle qui engendre w :

- si $N \rightarrow 11$: OK.
- si $N \rightarrow 1001$: OK.
- si $N \rightarrow N0$: alors $w = w_10$ avec $w_1 \in L$ donc par hypothèse de récurrence $b(w_1)$ est divisible par 3, or $b(w) = 2.b(w_1)$.
- si $N \rightarrow NN$: $w = w_1w_2$, $b(w) = b(w_1).2^k + b(w_2)$ est divisible par 3.

3. Est-ce que tous les nombres divisibles par 3 sont représentés dans L ?

Correction : Non, 21 s'écrit 10101 et n'est pas dans L .

Exercice 1.12

Quel sont les langages sur l'alphabet $V_t = \{ (,) \}$ engendrés par les grammaires qui suivent :

1. $S \rightarrow \varepsilon \mid (\mid) \mid SS$;
2. $S \rightarrow \varepsilon \mid (T \mid T \rightarrow S)$;
3. $S \rightarrow \varepsilon \mid (S) \mid SS$;
4. $S \rightarrow \varepsilon \mid (S)S$

Correction :

1. VT^* ;
2. $\{ (^n) \mid n \in \mathbb{N} \}$;
3. ensemble des mots bien parenthésés sur V_t ;
4. idem (cf. exercice 1.14).

Exercice 1.13

On considère la grammaire $G = (\{a, +, *\}, \{S\}, S, \{S \rightarrow SS * | SS + | a\})$. (C'est une version incomplète de la grammaire qui régit le fonctionnement d'une calculatrice HP.)

1. Trouvez les mots de L_G de longueur inférieure ou égale à 5 et donner une des dérivations qui vous permettent de les engendrer.

Correction :

- $\underline{S} \rightarrow a,$
- $\underline{S} \rightarrow \underline{SS} + \rightarrow a\underline{S} + \rightarrow aa +,$
- $\underline{S} \rightarrow \underline{SS} * \rightarrow a\underline{S} * \rightarrow aa *,$
- $\underline{S} \rightarrow \underline{SS} + \rightarrow \underline{SS} + S + \rightarrow a\underline{S} + S + \rightarrow aa + \underline{S} + \rightarrow aa + a +,$
- $\underline{S} \rightarrow \underline{SS} * \rightarrow \underline{SS} + S * \rightarrow a\underline{S} + S * \rightarrow aa + \underline{S} * \rightarrow aa + a *,$
- $\underline{S} \rightarrow \underline{SS} * \rightarrow \underline{SS} * S * \rightarrow a\underline{S} * S * \rightarrow aa * \underline{S} * \rightarrow aa * a *,$
- $\underline{S} \rightarrow \underline{SS} + \rightarrow \underline{SS} * S + \rightarrow a\underline{S} * S + \rightarrow aa * \underline{S} + \rightarrow aa * a +,$
- $\underline{S} \rightarrow \underline{SS} + \rightarrow a\underline{S} + \rightarrow a\underline{SS} + + \rightarrow aa\underline{S} + + \rightarrow aaa + +,$
- $\underline{S} \rightarrow \underline{SS} * \rightarrow a\underline{S} * \rightarrow a\underline{SS} + * \rightarrow aa\underline{S} + * \rightarrow aaa + *,$
- $\underline{S} \rightarrow \underline{SS} * \rightarrow a\underline{S} * \rightarrow a\underline{SS} * * \rightarrow aa\underline{S} * * \rightarrow aaa * *,$
- $\underline{S} \rightarrow \underline{SS} + \rightarrow a\underline{S} + \rightarrow a\underline{SS} * + \rightarrow aa\underline{S} * + \rightarrow aaa * +$

2. Dans la suite on note $|w|_\alpha$ le nombre d'occurrences de la lettre α dans le mot w .

Montrer par récurrence sur la longueur de la dérivation du mot que si $w \in L_G$ alors $|w|_+ + |w|_* + 1 = |w|_a$.

Correction : Pour le mots de longueur que l'on peut dériver en une étape, (ie a) c'est vrai. ($0 + 0 + 1 = 1$)

Les deux cas d'induction sont similaires. Faisons le cas $S \rightarrow SS +$.

Supposons que cela est pour tous les mots que l'on peut dériver en moins de n étapes. Soit w_1 et w_2 deux mots quelconques engendrés respectivement en $n-k$ et k ($n, k \geq 1$) étapes. On peut en n étapes engendrer les mots de la forme $w_1 w_2 +$. Par hypothèse d'induction, nous avons $|w_1|_+ + |w_1|_* + 1 = |w_1|_a$ et $|w_2|_+ + |w_2|_* + 1 = |w_2|_a$. Nous avons $|w_1 w_2 +|_+ = |w_1|_+ + |w_2|_+ + 1$, $|w_1 w_2 +|_* = |w_1|_* + |w_2|_*$ et $|w_1 w_2 +|_a = |w_1|_a + |w_2|_a$. Donc $|w_1 w_2 +|_+ + |w_1 w_2 +|_* + 1 = (|w_1|_+ + |w_2|_+ + 1) + (|w_1|_* + |w_2|_*) + 1 = (|w_1|_+ + |w_1|_* + 1) + (|w_2|_+ + |w_2|_* + 1) = |w_1|_a + |w_2|_a = |w_1 w_2 +|_a$.

3. Trouver une procédure qui permet de décider si un mot appartient au langage. On pourra utiliser une procédure auxiliaire qui maintient à jour une variable représentant le nombre d'expressions disponibles pour de futures opérations.

Correction : L'idée est de garder un trace à l'aide de la première variable du nombre d'arguments disponibles. En fin de mot, ce nombre d'argument doit être égal à un. Une opération "consomme" deux arguments et en rend un, le résultat de l'opération. De plus à aucun moment le nombre d'arguments disponibles ne doit être négatif. En CAML cela donne:

```
let decision =
  let rec est_dans_lg nb_arg = function
    "a"::tl -> est_dans_lg (nb_arg+1) tl
    |_ ::tl -> if nb_arg >= 2
                then est_dans_lg (nb_arg - 2 + 1) tl
                else false
    |_ -> nb_arg = 1
  in est_dans_lg 0;;
```

Exercice 1.14

Soit L_1 le langage engendré par la grammaire

$$S \rightarrow (S) \mid SS \mid \varepsilon$$

et L_2 le langage engendré par la grammaire

$$T \rightarrow (T)T \mid \varepsilon$$

1. Montrer que $L_2 \subseteq L_1$.

Correction : Si $w \in L_2$, alors soit $w = \varepsilon$ auquel cas $w \in L_1$; soit $w = (w_1)w_2$ avec $w_1, w_2 \in L_2$ par récurrence, on a $w_1, w_2 \in L_1$, mais alors w se dérive par la première grammaire : $S \rightarrow SS \rightarrow (S)S \rightarrow^* (w_1)w_2 = w$.

2. Montrer que L_2 est stable par concaténation, c'est-à-dire si $w_1 \in L_2$ et $w_2 \in L_2$ alors $w_1w_2 \in L_2$.

Correction : Si $w_1 = \varepsilon$, c'est clair, sinon $w_1 = (w_3)w_4$. Par récurrence on a $w_4w_2 \in L_2$ mais alors $T \rightarrow (T)T \rightarrow^* (w_3)w_4w_2$.

3. En déduire que $L_1 \subseteq L_2$.

Correction : Soit $w \in L_1$. Si $w = \varepsilon$: OK. Si $w = (w_1)$, $w_1 \in L_2$ par hypothèse de récurrence, $T \rightarrow (T)T \rightarrow (T) \rightarrow^* (w_1)$. Si $w = w_1w_2$: on conclut par l'hypothèse de récurrence et la propriété précédente.

Exercice 1.15

On considère la grammaire $G = (\{a, b\}, \{S, A, B\}, S, R)$ où R contient les règles

$$\begin{aligned} S &\rightarrow \varepsilon \mid aB \mid bA \\ A &\rightarrow aS \mid bAA \\ B &\rightarrow bS \mid aBB \end{aligned}$$

On note L_G le langage engendré par G .

1. Montrer que les mots $abab$ et $baaabb$ sont dans L_G .

Correction : $S \rightarrow aB \rightarrow abS \rightarrow abaB \rightarrow ababS \rightarrow abab$

$S \rightarrow bA \rightarrow baS \rightarrow baaB \rightarrow baaaBB \rightarrow baaabSB \rightarrow baaabB \rightarrow baaabbS \rightarrow baaabb$

2. Montrer que chaque $w \in L_G$ vérifie $|w|_a = |w|_b$. Indication : soit L_A l'ensemble de mots qui dérivent de A , et soit L_B l'ensemble de mots qui dérivent de B . Montrer que pour chaque mot w sur l'alphabet $\{a, b\}$, par récurrence sur la longueur de w , les trois assertions

- (a) si $w \in L_G$ alors $|w|_a = |w|_b$,
- (b) si $w \in L_A$ alors $|w|_a = |w|_b + 1$,
- (c) si $w \in L_B$ alors $|w|_a + 1 = |w|_b$.

Correction : Il faut faire une preuve par récurrence simultanée.

$w = \varepsilon$: $w \in L_G$, $w \notin L_A$, $w \notin L_B$.

$w = aw'$ et $w \in L_G$: on a $w' \in L_B$, alors par récurrence

$$|w|_a = |aw'|_a = 1 + |w'|_a = |w'|_b = |aw'|_b = |w|_b$$

$w = aw'$ et $w \in L_A$: on a $w' \in L_G$, alors par récurrence

$$|w|_a = |aw'|_a = 1 + |w'|_a = 1 + |w'|_b = 1 + |aw'| = 1 + |w|_b$$

$w = aw'$ et $w \in L_B$: on a $w' = vv'$ avec $v, v' \in L_B$, alors par récurrence

$$|w|_a + 1 = |avv'|_a + 1 = 1 + |v|_a + |v'|_a + 1 = |v|_b + |v'|_b = |avv'|_b = |w|_b$$

Les cas où w commence par b sont symétriques.

3. Montrer que pour chaque mot w sur l'alphabet $\{a, b\}$: si $|w|_a = |w|_b$ alors $w \in L_G$.
Indication: Montrer que pour chaque mot w sur l'alphabet $\{a, b\}$, par récurrence sur la longueur de w , les trois assertions

- (a) si $|w|_a = |w|_b$ alors $w \in L_G$,
- (b) si $|w|_a = |w|_b + 1$ alors $w \in L_A$,
- (c) si $|w|_a + 1 = |w|_b$ alors $w \in L_B$.

Correction : similaire



Exercice 1.16

Considérons le langage L des mots *non carrés* sur l'alphabet $\{a, b\}$, c'est-à-dire l'ensemble des mots w tels qu'il n'y ait pas de mot v tel que $w = vv$.

1. Donner une grammaire pour le langage des mots de longueur impaire de L .

Correction : On remarque que tous les mots de longueur impaire sont dans L .

$$A \rightarrow a \mid b \mid aaA \mid abA \mid baA \mid bbA$$

2. Montrer que si $w = v_1av_2bv_3$ avec $|v_2| = |v_1| + |v_3|$ alors $w \in L$.

Correction : Il faut montrer que w n'est pas un carré. Par l'absurde, si $w = vv$ alors $vv = v_1av_2bv_3$ mais comme $|v_2| = |v_1| + |v_3|$, on peut découper v_2 en v_4v_5 avec $|v_4| = |v_3|$ et $|v_5| = |v_1|$, et alors $|v_1av_4| = |v_5bv_3|$ donc $v = v_1av_4 = v_5bv_3$ mais cette dernière égalité est impossible car a et b sont à la même position.

3. En exprimant l'ensemble des mots de la forme ci-dessus (en découpant v_2 de manière astucieuse), donner une grammaire engendrant les mots de longueur paire de L .

Correction : On remarque que l'on peut également découper v_2 en v_4v_5 avec cette fois $|v_4| = |v_1|$ et $|v_5| = |v_3|$. Autrement l'ensemble des mots de la forme ci-dessus est

$$\{v_1av_4v_5bv_3 \mid |v_4| = |v_1| \text{ et } |v_5| = |v_3|\}$$

Ce langage est engendré par la grammaire

$$\begin{aligned} B &\rightarrow B_1B_2 \\ B_1 &\rightarrow a \mid aB_1a \mid aB_1b \mid bB_1a \mid bB_1b \\ B_2 &\rightarrow b \mid aB_2a \mid aB_2b \mid bB_2a \mid bB_2b \end{aligned}$$

On sait que tous les mots de ce langage sont dans L , réciproquement on remarque que les mots de L de longueur paire sont soit de la forme ci-dessus, soit de la forme obtenue en échangeant a et b . Par conséquent, une grammaire engendrant les mots de longueur paire de L est

$$B \rightarrow B_1B_2 \mid B_2B_1$$

4. En déduire une grammaire pour L .

Correction :

$$S \rightarrow A \mid B_1B_2 \mid B_2B_1$$

Chapitre 2

Grammaires régulières et automates

Ce chapitre a pour but de donner une méthode permettant de reconnaître les mots d'un langage. On s'intéresse à une classe de langages simples à reconnaître, les langages réguliers, reconnaissables par des automates.

2.1 Grammaires régulières

Définition 2.1 Une grammaire $G = (V_t, V_n, S, R)$ est dite régulière si toutes ses règles sont de la forme $N \rightarrow w$, avec $w \in V_t^*$ ou bien $w = uM$, $u \in V_t^*$ et $M \in V_n$. Un langage engendré par une grammaire régulière est dit régulier.

Les grammaires régulières sont beaucoup moins expressives que les grammaires hors-contexte. En particulier, elles ne peuvent pas engendrer le langage $\{a^n b^n \mid n \in \mathbb{N}\}$ (cf. exercice 2.3). Elles peuvent en fait engendrer des langages finis, ou des langages infinis dont la structure est simple.

Nous allons maintenant simplifier la forme de nos règles régulières en plusieurs étapes, sans en changer le pouvoir d'expression, c'est-à-dire que nous pourrions décrire les mêmes langages, avec bien sûr des règles différentes. Dans un premier temps, nous imposons que le membre droit contienne au plus un symbole terminal. Il reste quatre sortes de règles possibles :

$$N \rightarrow \varepsilon \qquad N \rightarrow \alpha \qquad N \rightarrow M \qquad N \rightarrow \alpha M$$

où $N \in V_n$, $\alpha \in V_t$ et $M \in V_n$. On peut en effet toujours transformer une règle de la forme $N \rightarrow \alpha_1 \cdots \alpha_n$ (resp. $N \rightarrow \alpha_1 \cdots \alpha_n M$) en n règles $N \rightarrow \alpha_1 N_1$, $N_1 \rightarrow \alpha_2 N_2$, $\dots$, $N_{n-1} \rightarrow \alpha_n$ (resp. $N_{n-1} \rightarrow \alpha_n M$) sans changer le langage engendré, au prix de l'ajout de $n-1$ nouveaux non-terminaux à la grammaire.

Par exemple, partant de la grammaire

$$\begin{aligned} S &\rightarrow ab \mid T \\ T &\rightarrow cdT \mid U \\ U &\rightarrow e \mid S \end{aligned}$$

on construit la grammaire

$$S \rightarrow aS_1 \mid T$$

$$\begin{aligned}
T &\rightarrow cT_1 \mid U \\
U &\rightarrow e \mid S \\
S_1 &\rightarrow b \\
T_1 &\rightarrow dT
\end{aligned}$$

On va maintenant éliminer les règles de la forme $N \rightarrow M$, en prenant soin d'ajouter les règles $N \rightarrow w$ pour toute règle $M \rightarrow w$ telle que $w \notin V_n$. La transformation à effectuer est en fait un peu plus complexe, il faut ajouter les règles en question pour tout non-terminal $M \neq N$ tel que $N \rightarrow^* M$, c'est-à-dire tel qu'il y ait une dérivation de N vers M .

Ainsi si on reprend la grammaire obtenue précédemment, on obtient

$$\begin{aligned}
S &\rightarrow aS_1 \mid cT_1 \mid e \\
T &\rightarrow cT_1 \mid aS_1 \mid e \\
U &\rightarrow e \mid aS_1 \mid cT_1 \\
S_1 &\rightarrow b \\
T_1 &\rightarrow dT
\end{aligned}$$

La grammaire obtenue à l'issue de ces deux phases ne nous satisfait pas encore entièrement car elle peut contenir des règles dont le membre droit est un seul terminal α . Les règles $N \rightarrow \alpha$ peuvent en fait être facilement éliminées à condition d'ajouter une règle $E \rightarrow \varepsilon$ (où E est un nouveau non-terminal) et de remplacer toute règle $N \rightarrow \alpha$ par $N \rightarrow \alpha E$.

Sur l'exemple précédent, cela donne

$$\begin{aligned}
S &\rightarrow aS_1 \mid cT_1 \mid eE \\
T &\rightarrow cT_1 \mid aS_1 \mid eE \\
U &\rightarrow eE \mid aS_1 \mid cT_1 \\
S_1 &\rightarrow bE \\
T_1 &\rightarrow dT \\
E &\rightarrow \varepsilon
\end{aligned}$$

Définition 2.2 Une grammaire régulière $G = (V_t, V_n, S, R)$ est dite réduite si toutes ses règles sont de la forme $N \rightarrow \alpha M$ ou $N \rightarrow \varepsilon$, où $N \in V_n$, $\alpha \in V_t$ et $M \in V_n$.

La discussion précédente montre que pour toute grammaire régulière, il existe une grammaire régulière réduite qui engendre le même langage.

Considérons la deuxième grammaire de la figure 1.1 engendrant les entiers en numération binaire, qui est régulière et dont la variante réduite est :

$$\begin{aligned}
N &\rightarrow 0E \mid 1M \\
M &\rightarrow 0M \mid 1M \mid \varepsilon \\
E &\rightarrow \varepsilon
\end{aligned}$$

Les dérivations de cette grammaire sont extrêmement simples, puisqu'il n'y a pas le choix du non-terminal à remplacer. Par exemple

$$N \rightarrow 1M \rightarrow 10M \rightarrow 100M \rightarrow 1001M \rightarrow 10011M \rightarrow 10011$$

De plus, il est facile de voir que cette grammaire est non-ambiguë, car la règle à utiliser est entièrement déterminée à chaque étape par le prochain symbole du mot à engendrer : il y a donc une unique dérivation pour chaque mot, et donc un unique arbre syntaxique (qui est filiforme). Cela n'est pas le cas de toutes les grammaires régulières réduites. Par exemple, la

première grammaire de la figure 1.1 engendrant les entiers en numération binaire, dont la variante réduite est :

$$\begin{aligned} N &\rightarrow 0E \mid 1E \mid 1M \\ M &\rightarrow \varepsilon \mid 0M \mid 1M \\ E &\rightarrow \varepsilon \end{aligned}$$

est ambiguë puisque le mot 1 a deux arbres syntaxiques différents. Nous verrons dans les paragraphes suivants un cas particulier de grammaire régulière réduite pour lequel la non-ambiguïté est garantie.

2.2 Automates déterministes

Nous allons maintenant nous préoccuper de reconnaître les mots engendrés par une grammaire régulière réduite. Le mécanisme est celui des automates.

Définition 2.3 *Un automate déterministe A est un quintuplet (V_t, Q, q_0, F, T) où*

1. V_t est le vocabulaire de A ;
2. Q est l'ensemble des états de A ;
3. $q_0 \in Q$ est l'état initial de A ;
4. $F \subseteq Q$ est l'ensemble des états acceptants de A ;
5. T est la fonction de transition, une fonction de $Q \times V_t$ dans Q .

Lorsque $T(q, \alpha) = q'$, on notera $q \xrightarrow{\alpha} q'$ ce qui se lit « dans l'état q , en lisant la lettre α on passe dans l'état q' ». Lorsque T est une application, autrement dit s'il y a dans chaque état exactement une transition pour toute lettre de l'alphabet, l'automate est dit *complet*.

On a l'habitude de dessiner les automates, en figurant les états par des cercles, en indiquant l'état initial par une flèche entrante, les états acceptants par un double cercle ou une flèche sortant d'un seul cercle, et la transition de l'état q à l'état q' en lisant la lettre α par une flèche allant de q vers q' et étiquetée par α . Les automates sont donc des graphes, dont les nœuds et les états sont étiquetés.

Reconnaître un mot par un automate fonctionne comme suit : les lettres du mot $\alpha_1 \dots \alpha_n$ sont entrées dans l'automate à partir de la gauche, c'est-à-dire dans l'ordre $\alpha_1, \dots, \alpha_n$. La lettre α_1 provoque une transition de l'état initial q_0 vers un état q_1 (qui peut être le même état que l'état initial), puis la lettre α_2 provoque une transition de l'état q_1 vers un état q_2 , etc. jusqu'à arriver à un état q_n . Il suffit alors de regarder si $q_n \in F$ auquel cas le mot est reconnu, alors qu'il est rejeté dans le cas contraire. Notons qu'il n'est pas possible d'être stoppé en chemin lorsque l'automate est complet. Si au contraire l'automate est incomplet, on convient que le mot n'est pas reconnu lorsque l'exécution est stoppée en chemin.

Définition 2.4 *Étant donné un automate $A = (V_t, Q, q_0, F, T)$, on appelle exécution toute suite (éventuellement vide) de transitions $q_0 \xrightarrow{\alpha_1} q_1 \dots q_{n-1} \xrightarrow{\alpha_n} q_n$, aussi notée $q_0 \xrightarrow{w}^* q_n$, qui lit le mot $w = \alpha_1 \dots \alpha_n$. On appelle alors q_n l'état atteint lors de l'exécution de A sur la donnée w .*

Un mot w est reconnu par l'automate A s'il existe une exécution de l'automate A qui se termine en ayant lu le mot w , et tel que l'état atteint soit acceptant. On note par $\text{Lang}(A)$ le langage des mots reconnus par l'automate A . Un langage reconnu par un automate est dit reconnaissable.

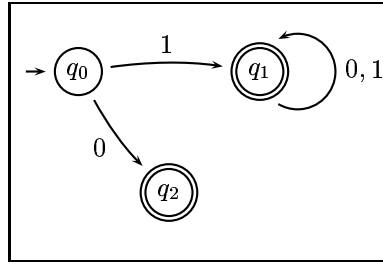


FIG. 2.1 – Automate déterministe reconnaissant les entiers naturels en numération binaire.

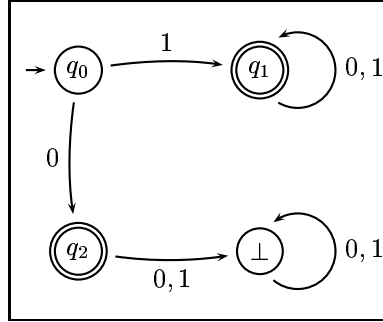


FIG. 2.2 – Automate déterministe complet reconnaissant les entiers naturels en numération binaire.

La figure 2.1 représente un automate reconnaissant le langage des entiers naturels en représentation binaire. Cet automate n'est pas complet (par exemple $T(q_2, 0)$ n'est pas défini), aussi il est susceptible de se bloquer lors d'une exécution, ce qui arrive si l'on veut exécuter cet automate sur le mot 00.

En pratique, il peut être utile de rendre un automate complet en ajoutant un nouvel état appelé *poubelle*, étiqueté par $\perp$, vers lequel vont toutes les transitions manquantes. L'automate obtenu reconnaît exactement le même langage si l'on convient que le nouvel état poubelle n'est pas acceptant. Ainsi, la figure 2.2 présente un automate complet reconnaissant le langage des entiers naturels en représentation binaire.

Notons que la reconnaissance d'un mot w par un automate est automatique : la lecture des lettres composant le mot provoque des transitions bien définies jusqu'à être bloqué (en cas de transitions manquantes) ou bien jusqu'à atteindre un état qui peut être ou ne pas être un état de satisfaction. Un automate déterministe est donc exempt d'ambiguïté.

2.3 Programmation des automates

2.3.1 Spécifications

La spécification du type abstrait « automate sur (V_t, Q) » est la suivante. Elle suppose donnée un alphabet V_t et un ensemble d'états Q .

$etat_initial$: automate sur $(V_t, Q) \rightarrow Q$
 $\{ etat_initial(A) \text{ est l'état initial de } A \}$
 $est_acceptant$: automate sur $(V_t, Q) \times Q \rightarrow \text{booléen}$

$\{ \text{est_acceptant}(A, q) \text{ est vrai si } q \text{ est acceptant dans } A \}$
 $\text{transition} : \text{automate sur } (V_t, Q) \times Q \times V_t \rightarrow Q$
 $\{ \text{transition}(A, q, \alpha) \text{ retourne l'état } T(q, \alpha) \text{ donné par } A \}$

Cette spécification exprime donc le fait qu'un automate sur un alphabet donné V_t et un ensemble d'états Q donné est caractérisé par les trois fonctions précédentes. Une implantation concrète de ce type de données dans un certain langage de programmation pourra être quelconque à condition de fournir au minimum ces trois fonctionnalités.

Les fonctions importantes sur les automates sont alors les suivantes.

fonction *exécute* : automate sur $(V_t, Q) \times Q \times V_t^* \rightarrow Q$
 $\{ \text{exécute}(A, q, w) \text{ retourne l'état final atteint par lecture du mot } w \text{ sur l'auto-}$
 $\text{mate } A \text{ à partir de l'état } q \}$
fonction *reconnaît* : automate sur $(V_t, Q) \times V_t^* \rightarrow \text{booléen}$
 $\{ \text{reconnaît}(A, w) \text{ retourne vrai si } w \in \mathcal{Lang}(A) \}$

2.3.2 Implantations

Il y a plusieurs façons d'implémenter ces spécifications dans un langage évolué, d'abord suivant le choix de ce langage, et ensuite suivant que l'on veuille implanter un automate particulier, ou bien implanter des fonctions génériques fonctionnant sur n'importe quel automate représenté par une structure de données.

2.3.3 Fonctions *exécute* et *reconnaît*

Mais pour commencer, regardons comment on peut implanter les fonctions *exécute* et *reconnaît*. Cette dernière se définit très simplement à partir des autres fonctions par

$$\text{reconnaît}(A, w) = \text{est_acceptant}(A, \text{exécute}(A, \text{etat_initial}(A), w))$$

se qui se traduit très facilement dans n'importe quel langage évolué.

La fonction *exécute* se traduit soit en une boucle (langage impératif), soit une fonction récursive (langage fonctionnel). Le principe est de toute façon le même : on parcourt les transitions jusqu'à ce que le mot d'entrée soit entièrement lu. Voici une version impérative de cette fonction :

```

fonction exécute(A:automate,q:etat,m:mot):etat
  tant que m n'est pas vide faire
    q = transition(A,q,tete(m))
    m = reste(m)
  fin tant que
  retourner q
fin fonction

```

et en voici une version fonctionnelle :

```

fonction exécute(A:automate,q:etat,m:mot):etat
  si m est vide
    q
  sinon
    soit q' = transition(A,q,tete(m))
    dans exécute(A,q',reste(m))
  fin si
fin fonction

```

2.3.4 Implantation du type automate

Il nous reste maintenant à montrer comment implanter les fonctionnalités des automates, c'est-à-dire principalement la fonction *transition*. Comme dit précédemment, cela dépend si l'on désire simplement coder un automate particulier, ou si l'on veut coder les automates en général.

Commençons par le premier cas qui est plus simple : il nous suffit d'écrire une fonction définie à l'aide de cas. Sur l'exemple de l'automate de la figure 2.2, cela donne la fonction suivante :

```

type etat =  $q_0, q_1, q_2, \perp$ 

fonction etat_initial():etat
  retourner  $q_0$ 
fin fonction

fonction est_acceptant(q:etat):etat
  si  $q=q_1$  ou  $q=q_2$  alors retourner vrai sinon retourner faux
fin fonction

fonction transition(q:etat, a:lettre):etat
  cas  $q=q_0$  :
    cas  $a=0$  : retourner  $q_2$ 
    cas  $a=1$  : retourner  $q_1$ 
  cas  $q=q_1$  : retourner  $q_1$ 
  cas  $q=q_2$  : retourner  $\perp$ 
  cas  $q=\perp$  : retourner  $\perp$ 
fin fonction

```

Dans le cas où l'on désire coder les automates en général, il s'agit de coder, une fois pour toutes, les fonctions *transition* et *est_acceptant* en fonction de la représentation choisie de l'automate. Comme d'habitude, ce choix de représentation est guidé par deux objectifs : l'efficacité des fonctions d'accès, ici *transition* et *est_acceptant*, et l'économie d'utilisation de la mémoire. Comme souvent, ces deux objectifs sont contradictoires, et la bonne représentation nécessite un compromis qui dépend de l'application.

Nous allons voir successivement deux représentations de la fonction de transition des automates. La première, la plus simple, consiste à représenter la fonction de transition sous la forme d'un tableau à deux dimensions. Sur notre exemple des entiers en binaire, cela donne :

	0	1
q_0	q_2	q_1
q_1	q_1	q_1
q_2	$\perp$	$\perp$
$\perp$	$\perp$	$\perp$

Pour faciliter la lecture du tableau, la première ligne indique les noms des lettres successives de l'alphabet de l'automate, et la première colonne les noms des états. La représentation mémoire omet bien sûr cette première ligne et cette première colonne, sachant que, *par convention*, les états sont numérotés de 1 à n , et les lettres de 1 à p . La transition de l'état de nom q_i en lisant la lettre de nom a_j est donc indiquée à la case (i, j) du tableau. Si le tableau a pour nom T , l'état obtenu est donc donné par $T(i, j)$. La fonction *transition* générique s'écrit donc

```

fonction transition(A:automate,q:etat,a:lettre):etat
  retourner  $T(q,a)$ 
fin fonction

```

On peut améliorer légèrement cette représentation en représentant un automate incomplet, ce qui se comprend car la dernière ligne est toujours remplie de $\perp$. Ceci donne sur notre exemple des entiers en numération binaire :

	0	1
q_0	q_2	q_1
q_1	q_1	q_1
q_2	$\perp$	$\perp$

Cela est facile à mettre en œuvre dans le calcul de la fonction *transition*:

```
fonction transition(A:automate,q:etat,a:lettre):etat
  si q= $\perp$ 
    retourner  $\perp$ 
  sinon
    retourner T(q,a)
  fin si
fin fonction
```

On peut aussi remarquer qu'une fois l'état poubelle atteint, le mot à lire ne peut pas être reconnu, et on peut donc accélérer la non-reconnaissance en retournant la valeur faux immédiatement. Cela nécessite bien sûr une modification de la fonction *exécute*.

Cette solution reste gourmande en mémoire, puisqu'elle occupe $n \times p$ cases mémoires. Cela est rédhibitoire pour les applications, par exemple à la preuve de circuits intégrés, qui nécessitent la manipulation d'automates pouvant avoir des millions d'états. On est alors amené à utiliser des techniques de représentation spécifiques (de type matrices creuses), en éliminant comme précédemment les transitions vers la poubelle, et en tirant parti de la structure de l'automate de manière à partager des parties communes. C'est ce qui est fait dans l'outil LEX que nous verrons au chapitre 3, les automates d'analyse lexicale pouvant être assez volumineux eux aussi (des centaines d'états et de lettres, donc des dizaines de milliers de cases dans le tableau).

La représentation permettant d'utiliser beaucoup moins de mémoire est les *listes de succession*. Cela consiste à associer à chaque état q la liste des couples (a, q') tels que $T(q, a) = q'$ pour lesquels q' n'est pas l'état poubelle. Cette représentation est proche du dessin de l'automate, et bien sûr beaucoup plus économe que la précédente lorsqu'il y a peu de transitions vers des états qui ne soient pas l'état poubelle. Sur l'exemple, cela donne :

q_0	$[(0, q_2); (1, q_1)]$
q_1	$[(0, q_1); (1, q_1)]$
q_2	$[\]$

Cette nouvelle représentation ne nécessite pas de numéroter les états, ni les lettres de l'alphabet, l'un et l'autre apparaissant dans la représentation. Elle est beaucoup plus économe en place puisqu'elle occupe $n \times 2t$ cases mémoires où t est le nombre moyen de transitions par états, en général petit devant p . Programmer la fonction *transition* est un peu plus complexe puisqu'il s'agit de rechercher la transition voulue dans la liste concernée :

```
fonction transition(A:automate,q:etat,a:lettre):etat
  l = liste_succession(A,q)
  retourner chercher(l,a)
fin fonction
```

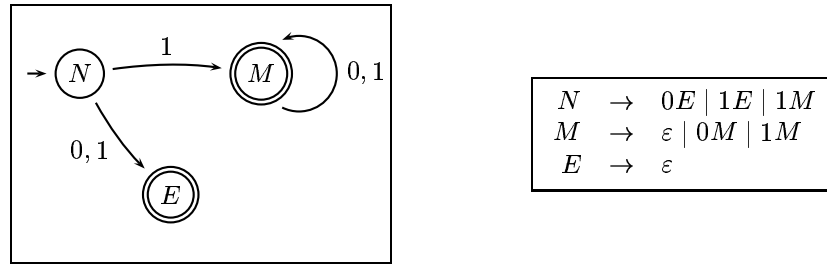


FIG. 2.3 – Automate et grammaires correspondantes pour les entiers naturels en notation binaire.

où la fonction *chercher* se code, en impératif, par

```

fonction chercher(l:liste,a:lettre):etat
  tant que l n'est pas vide
    (a',q')=tete(l)
    si a=a' retourner q'
    l = reste(l)
  fin tant que
  retourner ⊥
fin fonction

```

et en récursif par

```

fonction chercher(l:liste,a:lettre):etat
  si l est vide
    retourner ⊥
  sinon
    soit (a',q')=tete(l)
    si a=a'
      retourner q'
    sinon
      retourner chercher(reste(l),a)
    fin si
  fin si
fin fonction

```

2.4 Automates non-déterministes et grammaires régulières

L'ensemble des états d'un automate joue un rôle analogue aux non-terminaux d'une grammaire. Mais le vocabulaire des états n'a nul besoin d'être disjoint du vocabulaire terminal car ils ne sont jamais mélangés lors des exécutions. Pour comprendre la relation entre les deux notions, il suffit d'examiner un exemple de mot, par exemple le mot 10110, engendré par la grammaire et reconnu par l'automate de la figure 2.3.

On voit que la dérivation dans la première grammaire de la figure 2.3, et l'exécution de l'automate sont presque identiques :

dérivation :	$N \rightarrow 1M \rightarrow 10M \rightarrow 101M \rightarrow 1011M \xrightarrow{0} 10110E \rightarrow 10110$
exécution :	$N \xrightarrow{1} M \xrightarrow{0} M \xrightarrow{1} M \xrightarrow{1} M \xrightarrow{0} E$
mot reconnu :	1 10 101 1011 10110

De manière générale,

- à tout automate $A = (V_t, Q, q_0, F, T)$, on associe une grammaire régulière réduite $G = (V_t, Q, q_0, R)$ dont les non-terminaux sont les états de l'automate, l'état initial étant identifié avec la source, et dont les règles sont obtenues en ajoutant aux transitions de l'automate une règle produisant le mot vide pour chaque état de satisfaction :

$$R = \{q \rightarrow aq' \mid T(q, a) = q'\} \cup \{q \rightarrow \varepsilon \mid q \in F\}$$

- à toute grammaire régulière réduite $G = (V_t, V_n, S, R)$, on associe un automate $A = (V_t, V_n, S, F, T)$, où $F = \{M \mid M \rightarrow \varepsilon \in R\}$, et dont les transitions sont obtenues comme suit :

$$M \xrightarrow{a} N \text{ ssi } M \rightarrow aN \in R$$

La première transformation produit une grammaire régulière réduite qui est nécessairement non-ambiguë, puisqu'il n'y aura qu'une seule dérivation possible. Au contraire, si la grammaire de départ est ambiguë, la seconde transformation n'engendre pas nécessairement un automate déterministe. Partant de la grammaire

$$(\{0, 1\}, \{S, T\}, S, \{S \rightarrow 1S \mid 1T \mid \varepsilon, T \rightarrow 1T \mid \varepsilon\})$$

qui engendre le langage $\{1\}^*$, on obtient l'automate

$$(\{0, 1\}, \{S, T\}, S, \{S, T\}, \{S \xrightarrow{1} S, S \xrightarrow{1} T, T \xrightarrow{1} T\})$$

qui n'est pas un automate déterministe. En effet, sa fonction de transition n'est pas une application puisqu'elle associe à la donnée S le résultat S ou le résultat T . Cela montre que notre notation $q \xrightarrow{a} q'$ est plus générale que nécessaire, puisqu'elle permet de décrire des automates *non-déterministes* : ce sont des automates qui permettent plusieurs transitions depuis un état donné étiquetés par la même lettre, ce que l'on va formaliser en disant que la fonction de transition associe à chaque état et à chaque lettre un ensemble d'état.

Définition 2.5 Un automate non-déterministe A est un quintuplet (V_t, Q, q_0, F, T) où

1. V_t est le vocabulaire de A ,
2. Q est l'ensemble des états de A ,
3. $q_0 \in Q$ est l'état initial de A ,
4. $F \subseteq Q$ est l'ensemble des états acceptants de A ,
5. T , la fonction de transition de A , est une application de $Q \times V_t \rightarrow \mathcal{P}(Q)$.

La fonction de transition est maintenant à valeurs dans l'ensemble des parties de l'ensemble des états, $T(q, \alpha)$ est donc l'ensemble des états que l'on peut atteindre à partir de l'état q en lisant la lettre α . On continuera à noter $q \xrightarrow{\alpha} q'$ lorsque $q' \in T(q, \alpha)$.

Un automate déterministe est alors un cas particulier d'automate non-déterministe qui associe à tout élément de $Q \times V_t$ une partie de Q possédant zéro ou un élément (exactement un si c'est un automate complet). La reconnaissance d'un mot par un automate non-déterministe implique la *recherche* d'une exécution qui réussisse parmi toutes les exécutions possibles :

Définition 2.6 Étant donné un automate non-déterministe $A = (V_t, Q, q_0, F, T)$, on appelle exécution associée au mot $\alpha_1 \dots \alpha_n \in V_t^*$ toute suite (éventuellement vide) de transitions $q_0 \xrightarrow{\alpha_1} q_1 \dots q_{n-1} \xrightarrow{\alpha_n} q_n$, qui lit le mot $\alpha_1 \times \dots \times \alpha_n$. Un mot w est reconnu par l'automate A s'il existe une exécution de l'automate A qui se termine en état acceptant en ayant lu le mot w . On note par $\text{Lang}(A)$ le langage des mots reconnus par l'automate A .

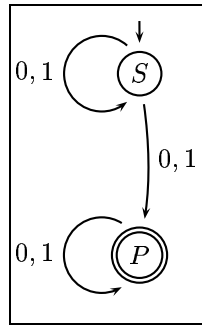


FIG. 2.4 – Un automate non-déterministe.

Par exemple, considérons l'automate de la figure 2.4 qui reconnaît le langage $\{0, 1\}^+$, et le mot 10110 de ce langage. Des deux exécutions qui suivent, seule la seconde est un succès

$$S \xrightarrow{1} S \xrightarrow{0} S \xrightarrow{1} S \xrightarrow{1} S \xrightarrow{0} S, \quad S \xrightarrow{1} S \xrightarrow{0} S \xrightarrow{1} S \xrightarrow{1} S \xrightarrow{0} P$$

mais cela suffit pour que le mot 10110 soit reconnu.

2.5 Déterminisation des automates

On pourrait croire que la classe des automates non-déterministes est strictement plus expressive que la classe des automates déterministes. En fait, il n'en est rien : il est toujours possible de *déterminiser* un automate en ajoutant de nouveaux états : si Q est l'ensemble des états d'un automate non-déterministe, l'ensemble des états de l'automate déterministe associé sera $\mathcal{P}(Q)$, l'ensemble des parties de Q . L'idée est que s'il est possible dans l'automate non-déterministe d'atteindre les états $q_1, \dots, q_n$ depuis l'état q en lisant la lettre a , alors il sera possible dans l'automate déterministe d'atteindre l'état $\{q_1, \dots, q_n\}$ en partant de tout état contenant q et en lisant la lettre a .

Définition 2.7 Soit $A = (V_t, Q, q_0, F, T)$ un automate non-déterministe. On définit $\text{Det}(A)$ comme l'automate $(V_t, \mathcal{P}(Q), \{q_0\}, F_{\text{det}}, T_{\text{det}})$ où

$$\begin{aligned} F_{\text{det}} &= \{K \in \mathcal{P}(Q) \mid K \cap F \neq \emptyset\} \\ T_{\text{det}}(K, a) &= \bigcup_{q \in K} T(q, a) \end{aligned}$$

Théorème 2.8 Soit A un automate non-déterministe. Alors $\text{Det}(A)$ est déterministe et reconnaît le même langage que A .

Par exemple, l'automate de la figure 2.5 est le déterminisé de celui de la figure 2.4. Notons que l'automate déterminisé obtenu à partir d'un automate non-déterministe possède en général un grand nombre d'états inutiles car inaccessibles :

Définition 2.9 L'état $q \in Q$ est accessible dans l'automate $A = (V_t, Q, q_0, F, T)$ s'il existe un mot $w \in V_t^*$ tel que $q_0 \xrightarrow{w}^* q$.

Les états inaccessibles peuvent être facilement éliminés en calculant les états accessibles depuis l'état initial. Dans l'automate déterminisé de la figure 2.5, les états $\{S\}$ et $\{S, P\}$ sont accessibles, et $\emptyset$ et $\{P\}$ sont inaccessibles.

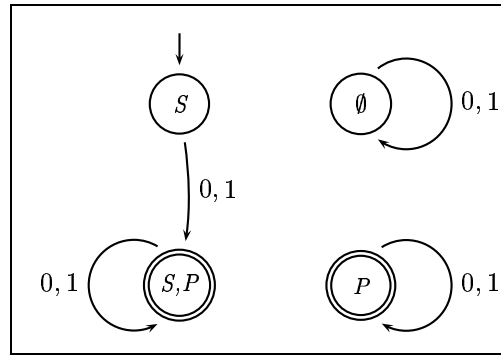


FIG. 2.5 – Automate déterminisé.

2.6 Minimisation des automates

Dans tout ce paragraphe, il est fondamental de supposer les automates COMPLETS, même si leur représentation en graphes correspond parfois à un automate incomplet par souci de clarté des dessins. On supposera également que tous les états des automates manipulés sont accessibles.

On voit que la détermination fait exploser le nombre d'états, mais que tous les états de $\mathcal{P}(Q)$ ne sont pas utiles à la construction du déterminisé d'un automate dont Q est l'ensemble des états. De plus, il se trouve que l'automate construit est en fait souvent plus compliqué que nécessaire. Un premier exemple est l'automate déterminisé de la figure 2.5 où l'on peut obtenir un automate plus simple en supprimant les états inaccessibles. L'élimination des états inaccessibles n'est pas toujours suffisante pour obtenir l'automate le plus simple possible. En haut de la figure ?? sont représentés trois automates équivalents reconnaissent le même langage $\{0(01)^n00 \mid n \geq 0\} \cup \{1(01)^n00 \mid n \geq 0\}$ (la vérification, facile au demeurant, que ces trois automates reconnaissent bien ce langage est laissée au lecteur). Le premier de ces trois automates comporte moins d'états, il est en fait minimal. Les deux automates (non-minimaux) appelés AutomateG1 et AutomateG2 ont en fait été obtenus par une déformation d'Automate initial. Les autres automates apparaissant sur la figure correspondent à des transformations réduisant le nombre d'états que nous allons maintenant expliquer.

Il est clair que les états finaux de l'automate AutomateG1, ainsi que ceux de l'automate AutomateD1 sont interchangeables. Étant à l'un quelconque de ces états, on est en état final, donc on reconnaît le mot vide, et toute transition nous envoie dans la poubelle (non représentée sur le dessin). Plus précisément, ils reconnaissent le même langage $\{\varepsilon\}$. On va donc les confondre, c'est ce qui est fait à cette première étape, qui nous rend donc les automates AutomateG2 et AutomateD2.

D'une manière générale, si deux états d'un automate reconnaissent le même langage, on les dira *interchangeables* et on peut alors *réduire* l'automate. Notons par $\mathcal{L}_A^q$ le langage des mots reconnus par l'automate obtenu à partir de $\mathcal{A}$ en adoptant l'état q comme état initial. Dans l'exemple précédent, $\mathcal{L}_{\text{AutomateG1}}^4 = \mathcal{L}_{\text{AutomateG1}}^7 = \{\varepsilon\}$. Le problème, bien sûr, va être de déterminer si deux états q et q' sont interchangeables. La réponse à cette question est difficile, et nous nous contenterons d'approximer la notion d'interchangeabilité par la relation d'équivalence $\equiv$ ainsi définie:

$$q \equiv q' \text{ ssi } (q \in F \text{ mboxssi } q' \in F) \text{ and } \forall a \in V_t \ T(q, a) = T(q', a)$$

Il est aisé de vérifier que la relation ainsi définie est réflexive ($q \equiv q$), symétrique ($q \equiv q'$ implique $q' \equiv q$) et transitive ($q \equiv q'$ et $q' \equiv q''$ implique $q \equiv q''$), et que deux états équivalents sont interchangeables, c'est-à-dire reconnaissent le même langage.

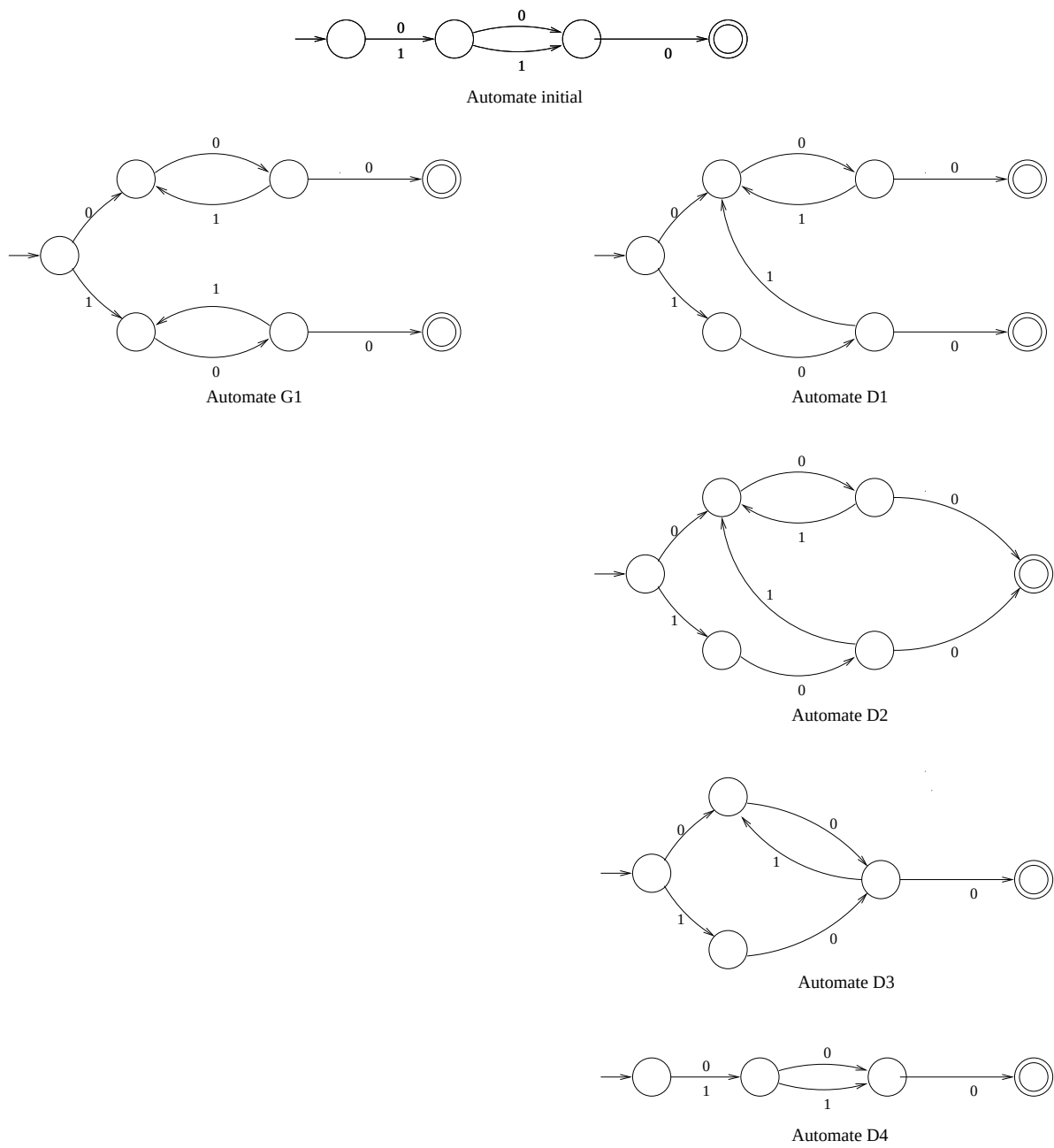


FIG. 2.6 – Réduction d'un automate non-minimal.

Supposons que les états $q_1, \dots, q_p$ de l'automate $\mathcal{A}$ soient équivalents. On va construire l'automate réduit $\mathcal{A}_{\text{réduit}} = (V_t, Q \setminus \{q_2, \dots, q_p\}, q_0, T_{\text{réduit}})$ ainsi défini :

- (i) q_1 est initial si l'un des états $q_1, \dots, q_p$ est initial,
- (ii) $T_{\text{réduit}}(q, a) = T(q, a)$ si $q \notin \{q_1, \dots, q_p\}$ et $T(q, a) \notin \{q_1, \dots, q_p\}$
 $T_{\text{réduit}}(q, a) = q_1$ si $T(q, a) \in \{q_1, \dots, q_p\}$

Notons que les arcs sortant d'un état q_i étant redondants dans l'automate de départ avec ceux sortant de l'état q_1 , il n'est nul besoin de les recopier dans l'automate résultat.

Il est facile de vérifier que cet automate est lui-aussi déterministe (aucun arc sortant n'ayant été ajouté) et qu'il reconnaît le même langage que l'automate de départ. En effet, tout mot reconnu dans l'automate $\mathcal{A}$ sans passer par un état de $q_1, \dots, q_n$ sera encore reconnu de la même manière. Un mot reconnu par un chemin empruntant un arc $q \xrightarrow{a} q_i$, avec $q \notin \{q_1, \dots, q_p\}$ sera maintenant reconnu avec le changement de cet arc en l'arc $q \xrightarrow{a} q_1$, la suite de la reconnaissance se poursuivant normalement à partir de l'état q_1 puisque les arcs sortants de q_1 et de q_i sont les mêmes par hypothèse.

Dans notre cas, il est facile de réduire l'automate **AutomateD1** en l'automate initial, grâce à 3 réductions successives. Notons que chaque réduction (sauf la dernière) fait apparaître de nouvelles réductions. Par contre, une seule réduction s'applique à l'automate **AutomateG1**, et l'automate minimal ne peut donc pas être atteint. Cela montre les limites de la méthode, lorsque nous détectons des états interchangeable grâce à l'équivalence $\equiv$.

Le but de la réduction précédente est de partitionner l'ensemble des états en états équivalents, c'est-à-dire reconnaissant le même langage. La partition obtenue a la propriété clé de *compatibilité* :

Définition 2.10 *Étant donné un automate déterministe complet $\mathcal{A} = (V_t, Q, q_0, F, T)$, une partition $\mathcal{P}$ de Q est compatible si toute classe de la partition est compatible. Une classe C est compatible si elle vérifie la propriété suivante :*

$$\forall q, q' \in C \quad \forall a \in V_t \quad \exists C', T(q, a) \in C' \text{ ssi } T(q', a) \in C'$$

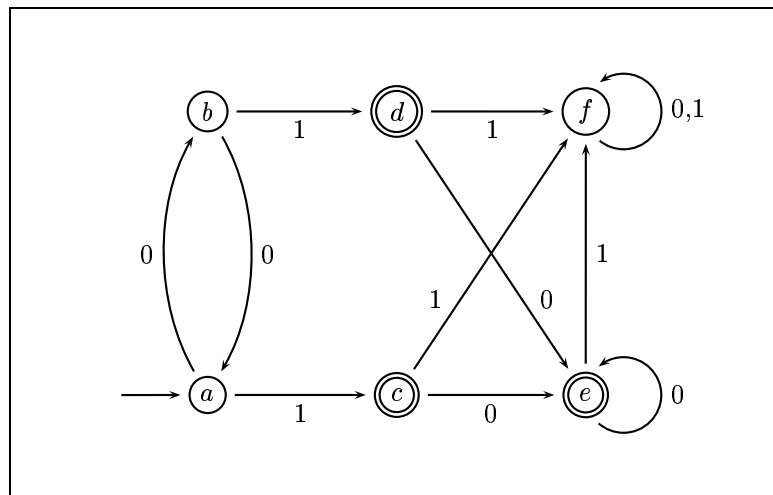
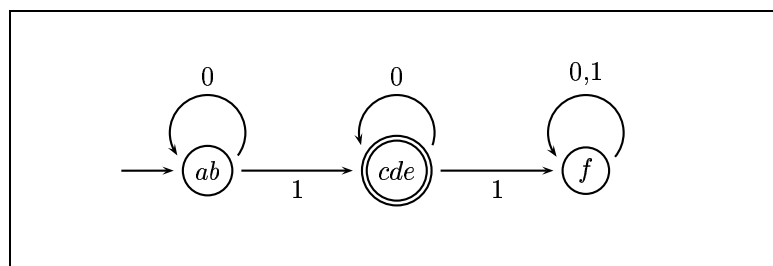
Étant donné une partition compatible des états Q de l'automate $\mathcal{A}$, on peut aisément construire un nouvel automate qui reconnaît le même langage que $\mathcal{A}$:

Proposition 2.11 *Étant donnée une partition compatible $\mathcal{P}$ des états de l'automate $\mathcal{A} = (V_t, Q, q_0, F, T)$, l'automate quotient $\mathcal{A}_q = (V_t, \mathcal{P}, C_0, \{C \in \mathcal{P} \mid C \subseteq F\}, T_q)$ tel que $T_q(C, a) = C'$ ssi $\exists q \in C \quad \exists q' \in C'$ tel que $T(q, a) = q'$ reconnaît le $\mathcal{L}(\mathcal{A})$.*

Notons que la compatibilité de la partition assure que la classe C' ne dépend pas du choix de l'état q de C , et que de plus, C' existe toujours puisque $\mathcal{A}$ est complet.

Nous sommes donc ramenés à trouver la partition compatible de Q dont le nombre de classes est minimum. Le problème est que la détermination a-priori d'une telle partition n'est pas simple, et peut difficilement être obtenue par un algorithme de réduction comme le précédent. Nous allons au contraire procéder à l'envers, c'est ce que nous expliquons maintenant.

L'idée est de partir d'une partition des états en deux classes, les états acceptants d'une part, et les autres, puisqu'une classe compatible ne peut contenir à la fois des acceptants et des non-acceptants. On va ensuite affiner cette partition chaque fois que nécessaire, c'est-à-dire lorsque une classe C ne sera pas compatible. Ce sera le cas lorsqu'une même lettre a provoque la transition de $p \in C$ vers p' et de $q \in C$ vers q' , p' et q' n'appartenant pas à une même classe, donc ne reconnaissant pas le même langage. On doit alors partitionner C en sous classes compatibles, ce qui fera apparaître au moins deux nouvelles classes, C' contenant p et C'' contenant q . Ce partitionnement va rendre incompatibles des classes auparavant compatibles, et il va donc falloir tester toutes les classes à l'exception de celles

FIG. 2.7 – *Un automate non minimal.*FIG. 2.8 – *Automate minimisé.*

issues du partitionnement de C . L'algorithme s'arrête (nécessairement) lorsque toutes les classes sont compatibles.

Étant donnée une partition P des états d'un automate, la fonction suivante recherche une classe de P qui n'est pas compatible. Dans le cas où toutes les classes sont compatibles la fonction retourne $\emptyset$. Puisque tous les classes dans les partitions construites par l'algorithme sont non-vides, le résultat $\emptyset$ signale le cas d'échec de la recherche.

```

type classe sur  $Q$  = ensemble sur  $Q$ 
type partition sur  $Q$  = ensemble sur classe sur  $Q$ 
fonction non_compatible( $A$ : automate sur  $(V_t, Q)$ ,
                         $P$ : partition sur  $Q$ ): classe sur  $Q$ 
% Recherche d'une classe non compatible de  $P$ 
si il existe  $C \in P$ ,  $p \in C$ ,  $q \in C$  et  $a \in V_t$  tels que
     $transition(A, p, a) \in C'$  et  $transition(A, q, a) \in C''$  avec  $C' \neq C''$ 
    retourner  $C$ 
sinon
    retourner  $\emptyset$ 
fin si
fin fonction

```

La fonction suivante calcule alors, à partir d'une partition P donnée, le *raffinement compatible le plus grossier* de P , c'est-à-dire la sous-partition compatible de P ayant le plus petit nombre possible de classes (on peut montrer qu'elle est unique) :

```

fonction partition( $A$ : automate sur  $(V_t, Q)$ ,
                     $P$ : partition sur  $Q$ ): partition sur  $Q$ 
% Calcul d'un raffinement compatible de la partition  $P$ 
soit  $C = non\_compatible(A, P)$ 
si  $C = \emptyset$ 
    retourner  $P$ 
sinon
    soit  $P'$  la partition la plus grossière de  $C$  telle que
        pour tout  $C' \in P'$ , pour tout  $C'' \in P$ , pour tous  $p, q \in C'$ 
        pour tout  $a \in V_t$  :
             $transition(A, p, a) \in C''$  ssi  $transition(A, q, a) \in C''$ 
    retourner partition( $A, (P - \{C\}) \cup P'$ )
fin si
fin fonction

```

La fonction suivante calcule l'automate minimisé.

```

fonction minimisation( $A$ : automate sur  $(V_t, Q)$ ): automate sur  $(V_t, \text{classe sur } Q)$ 
% Calcul du minimisé de  $A$ 
% Calcul de la plus grossière partition compatible
soit  $F = \{q \in Q \mid est\_acceptant(A, q)\}$ 
 $P = partition(A, \{Q - F, F\})$ 
% Calcul de l'état initial de l'automate minimal
soit  $I_{min}$  la classe de  $P$  telle que  $etat\_initial(A) \in I_{min}$ 
% Calcul des états acceptants de l'automate minimal
pour chaque  $C \in P$  :
    soit  $F_{min}(C) = true$  ssi  $C \cap F \neq \emptyset$ 
fin pour
% Calcul de la fonction de transition de l'automate minimal

```

```

pour chaque  $(C, a) \in P \times V_t$  faire
  soit  $q \in C$  (un représentant quelconque)
  soit  $C'$  la classe de  $P$  telle que  $\text{transition}(A, q, a) \in C'$ 
   $T_{\min}(C, a) = C'$ 
fin pour
retourner  $(I_{\min}, F_{\min}, T_{\min})$ 
fin fonction

```

Appliqué à l'exemple précédent, cet algorithme calcule la séquence de partitions

$$\frac{\{a, b, f\}, \{c, d, e\}}{\{a, b\}, \{f\}, \{c, d, e\}}$$

et retourne l'automate de la figure 2.8.

Théorème 2.12 *Soit $A = (V_T, Q, q_0, F, T)$ un automate déterministe et complet. L'automate envoyé par la fonction minimisation est l'automate déterministe complet avec un nombre minimal d'états acceptant le même langage que A .*

2.7 Exercices

Exercice 2.1

1. Donner une grammaire régulière décrivant les réels en virgule flottante, avec un signe éventuel.

Correction : Pour une grammaire régulière, les membres droits doivent être w ou wN où $w \in V_t^*$ et $N \in V_n$. On peut prendre :

$$\begin{aligned}
 \text{Réal} &\rightarrow + \text{Réal_sans_signe} \mid - \text{Réal_sans_signe} \mid \text{Réal_sans_signe} \\
 \text{Réal_sans_signe} &\rightarrow 0 \text{Réal_sans_signe} \mid \dots \mid 9 \text{Réal_sans_signe} \mid 0. \text{Entier} \mid \dots \mid 9. \text{Entier} \\
 \text{Entier} &\rightarrow 0 \text{Entier} \mid \dots \mid 9 \text{Entier} \mid 0 \mid \dots \mid 9
 \end{aligned}$$

2. Donner une grammaire régulière réduite équivalente.

Correction : Les membres droits doivent être ε ou bien αN avec $\alpha \in V_t$.

$$\begin{aligned}
 \text{Réal} &\rightarrow + \text{Réal_sans_signe} \mid - \text{Réal_sans_signe} \mid \\
 &\quad 0 \text{Suite_réel_sans_signe} \mid \dots \mid 9 \text{Suite_réel_sans_signe} \\
 \text{Réal_sans_signe} &\rightarrow 0 \text{Suite_réel_sans_signe} \mid \dots \mid 9 \text{Suite_réel_sans_signe} \\
 \text{Suite_réel_sans_signe} &\rightarrow 0 \text{Suite_réel_sans_signe} \mid \dots \mid 9 \text{Suite_réel_sans_signe} \mid . \text{Entier} \\
 \text{Entier} &\rightarrow 0 \text{Suite_entier} \mid \dots \mid 9 \text{Suite_entier} \\
 \text{Suite_entier} &\rightarrow 0 \text{Suite_entier} \mid \dots \mid 9 \text{Suite_entier} \mid \varepsilon
 \end{aligned}$$

3. Dessiner un automate reconnaissant les réels.

Exercice 2.2

Transformer la grammaire suivante en une grammaire régulière réduite:

$$S \rightarrow a \mid T \mid \varepsilon \qquad T \rightarrow bcT \mid U \qquad U \rightarrow e \mid S \mid \varepsilon$$

**Exercice 2.3**

On veut montrer que le langage $L = \{a^n b^n \mid n \in \mathbb{N}\}$ ne peut pas être engendré par une grammaire régulière.

1. Soit $A = (V_t, Q, q_0, F, T)$ un automate déterministe quelconque. Trouver une borne $N \in \mathbb{N}$ associée à A qui assure que pour tout mot $w \in \text{Lang}(A)$ de longueur supérieure ou égale à N , l'exécution de A sur w passe nécessairement au moins deux fois par un certain état.
2. Raisonnons par l'absurde en supposant que le langage L est reconnaissable par un automate A . Par un raisonnement analogue à la question précédente, trouver un entier N tel que lors de l'exécution de A sur le mot $a^N b^N$, un état est parcouru au moins deux fois lors de la lecture des lettres a .
3. En déduire qu'il existe un mot de la forme $a^{N+k} b^N$ avec $k > 0$ reconnu par l'automate, et conclure.

Exercice 2.4

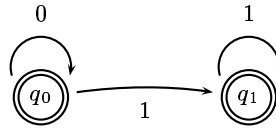
Donner un automate qui reconnaisse l'ensemble des entiers naturels en représentation octale.

Exercice 2.5

Pour chacun des langages suivants sur l'alphabet $\{0, 1\}$, trouver un automate qui le reconnaisse.

1. $\{0^n 1^m \mid n, m \geq 0\}$

Correction :

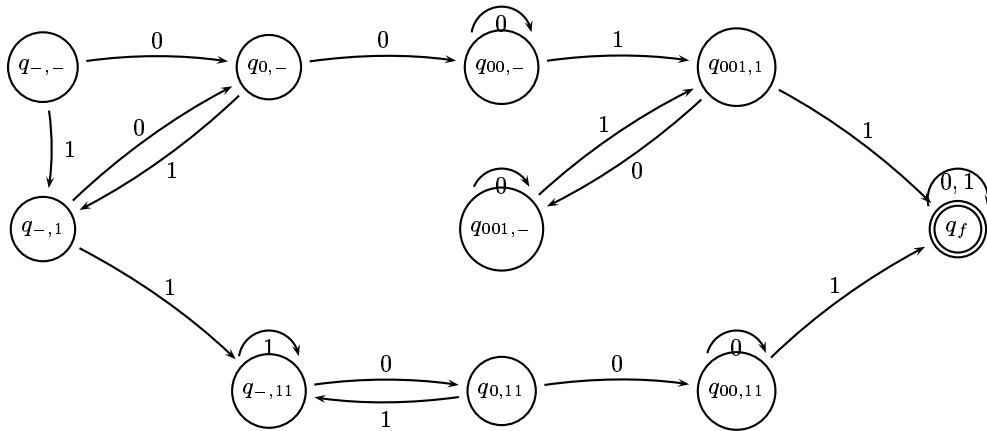


2. l'ensemble des mots contenant 11,
3. l'ensemble des mots contenant 001,
4. l'ensemble des mots contenant 001 et 11,

Correction :

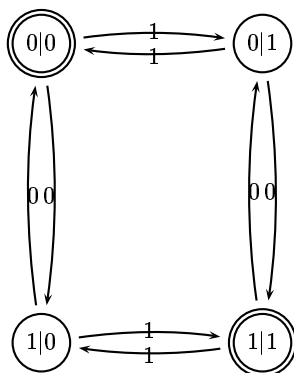
On peut indiquer les états par 2 composantes qui disent si on a déjà rencontré l'une des séquences, et qui se rappellent les 2 ou 3 dernières lettres. Remarquer que le mot 0011 convient.

A vérifier



5. l'ensemble des mots contenant 001 ou 11.
6. L'ensemble des mots tel que la parité du nombre de 0 est égale à celle de 1.

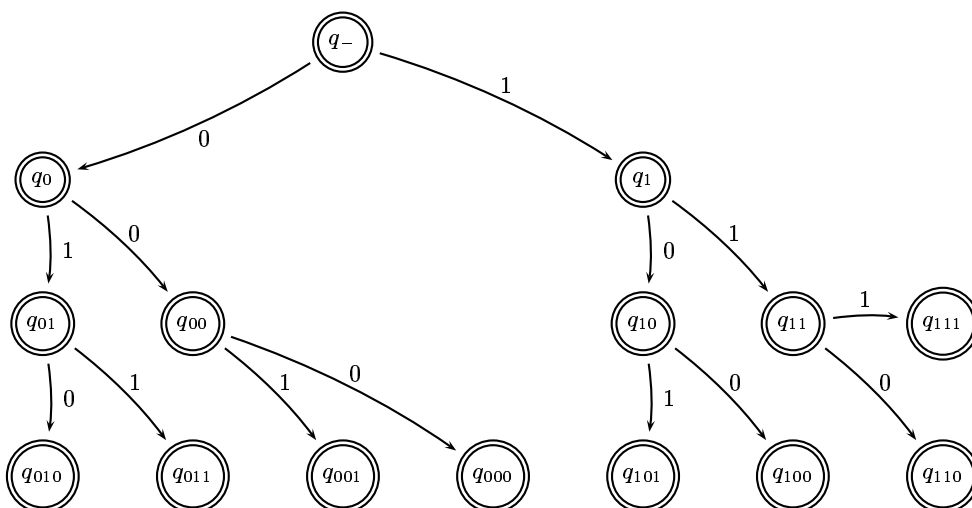
Correction :



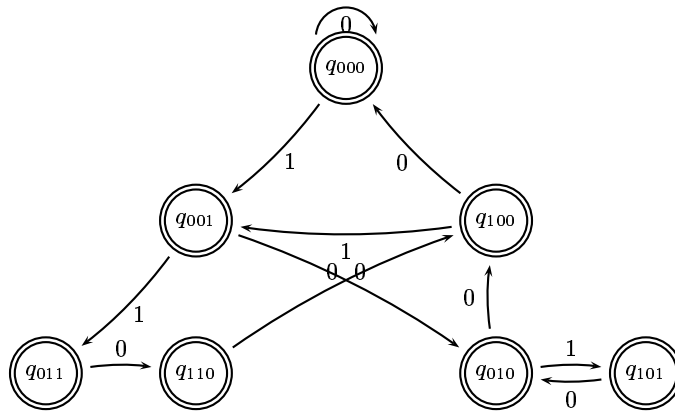
7. l'ensemble des mots où chaque bloc de 4 caractères successifs contient au moins deux fois 0,

Correction :

En particulier il faut accepter tout les mots de moins de 4 lettres. L'automate est en deux parties, une qui atteint tous les états de trois lettres, l'autre est un cycle reliant ces états a trois lettres. Il suffit de se rapeller de trois lettres, car la quatrieme est lue.

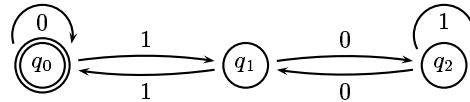


Avec :



8. l'ensemble des nombres naturels en représentation binaire qui sont divisibles par 3.

Correction : On mémorise le reste de la division par 3 dans un état, sachant que la lecture d'un 0 fait multiplier par 2, et la lecture d'un 1 fait $*2 + 1$. Remarque: pour simplifier 000... est un nombre binaire accepté. On pourra montrer la transformation pour avoir la solution stricte.



Exercice 2.6

On appelle *chien de garde* un système qui reçoit les signaux $\{e, f, r, s\}$ et qui est spécifié ainsi :

- Initialement le système est en état de repos et ignore les messages recus.
- À l'arrivée du signal e (éveil), le système passe à l'état de veille.
- Quand le système est en veille, tout signal s le fait passer dans un état d'alerte.
- Le signal f (fin) le fait toujours passer dans l'état de repos.
- Le système reste en état d'alerte tant que le signal r (repos) n'est pas émis.

Donner un automate réalisant un tel chien de garde.

Correction : Une introduction. Implicitement l'état final est celui de repos.

Exercice 2.7

1. Donner un automate déterministe qui reconnaît une date de la forme "jour/mois" où le jour est noté par un entier entre 1 et 31 et le mois est noté par un entier entre 1 et 12. Attention aux mois avec moins de 31 jours! Par exemple, votre automate doit accepter "7/4" (pour le 7 avril) et "31/12" (pour le 31 décembre), mais pas "30/2".
2. Étendre l'exercice aux dates de la forme "jour/mois/année", où l'année est notée par un entier entre 1000 et 1999. On suppose que les années bissextiles sont les années qui sont divisibles par 4.

Exercice 2.8

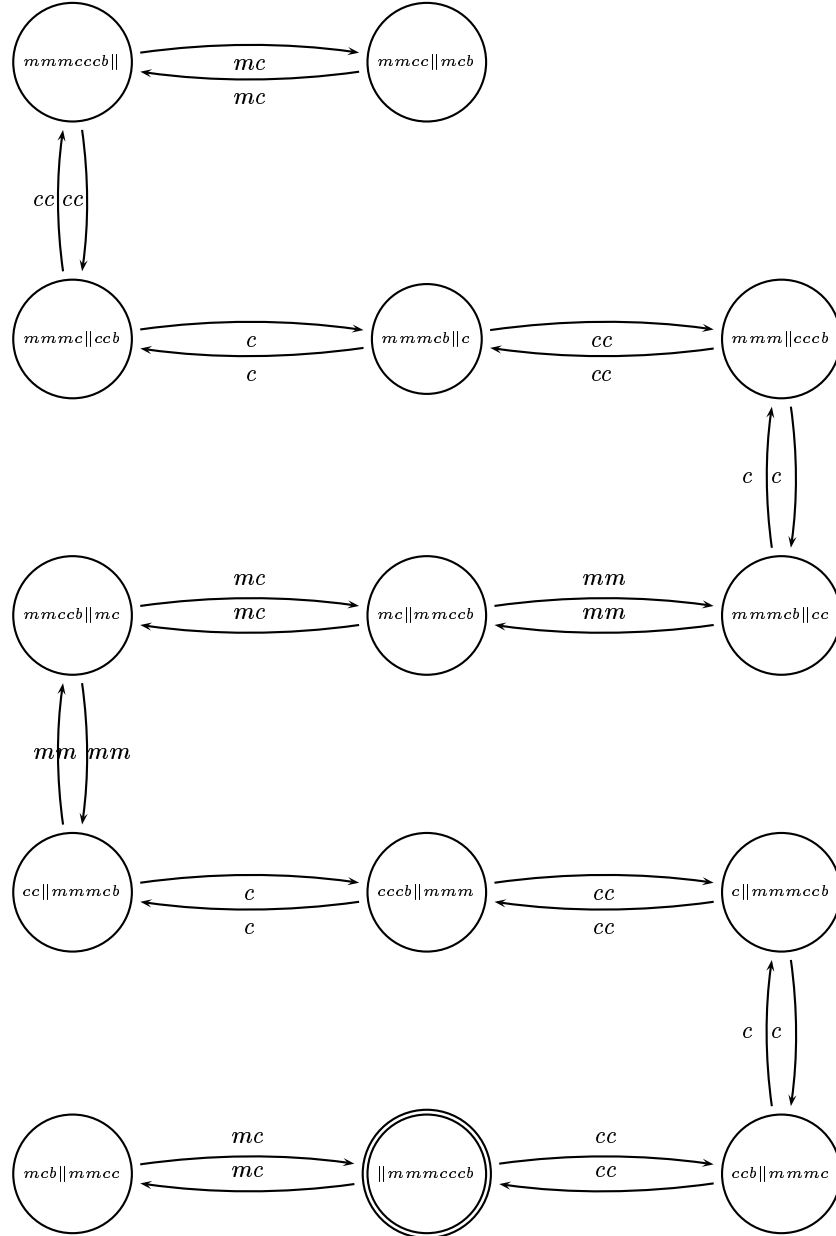
Trois missionnaires et trois cannibales veulent traverser une rivière pleine de piranhas. Il y a une seule pirogue qui peut transporter deux personnes au maximum. Un missionnaire ne

peut pas utiliser la pirogue tout seul (mais deux le peuvent). Dès qu'il y a, à quelque endroit, plus de cannibales que de missionnaires, les cannibales vont manger les missionnaires.

- Est-ce qu'il y a une possibilité pour les six personnes de traverser la rivière sans être mangées?

Indication : utiliser un automate où les états sont les différentes répartitions possibles des missionnaires et des cannibales d'un côté et de l'autre de la rivière, et les transitions représentent les traversées de la pirogue.

Correction :



Un missionnaire est figuré par un M, un cannibale par un C et le bateau par un B.

- Si l'on considère le mot des actions, que penser d'une solution miroir? Justifier.

Correction : C'est encore une solution, traverser d'un coté pour aller à l'autre c'est la même chose que le contraire. Reste à le justifier.

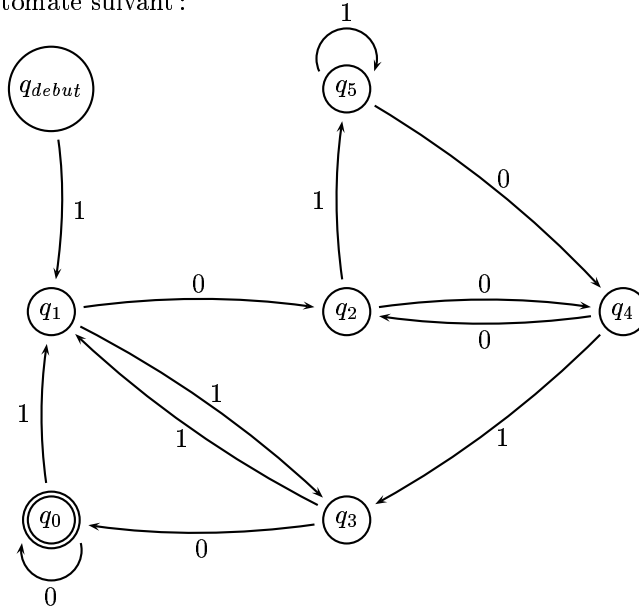
Pour avoir l'ensemble des solutions, on a un automate complet, et localement réversible ($q \xrightarrow{a} q' \Rightarrow q' \xrightarrow{a} q$) Soit $\tilde{q}_{b\|a} = q_{a\|b}$ pour chaque $q_{a\|b}$ de notre automate.

On montre par récurrence sur $\|w\|$, mot quelconque, si $q_i \xrightarrow{w} q_j \Rightarrow \tilde{q}_j \xrightarrow{\text{mir}(w)} \tilde{q}_i$

Dans le cas où $w = w_1.a$, si $q_i \xrightarrow{w_1.a} q_j$ alors $q_i \xrightarrow{w_1} q_k \xrightarrow{a} q_j$ et par récurrence $\tilde{q}_j \xrightarrow{a} \tilde{q}_k \xrightarrow{\text{mir}(w_1)} \tilde{q}_i$

Exercice 2.9

On considère l'automate suivant :



1. Vérifiez que les mots 110, 1100, 10010 sont reconnus par cet automate.
2. Caractériser le langage L_0 reconnu par cet automate?

Correction : C'est l'ensemble des mots non nuls dont la valeur en binaire est multiple de 6. Si $b(w)$ est la valeur de w en binaire, on montre que si $q_{debut} \xrightarrow{w} q_1$ alors le reste de la division de $val(w)$ par 6 vaut 1. Le cas d'une lettre est facile, si $w = w_1.a$, alors puisque l'automate est déterministe et complet, il existe q_i et q_j tels que $q_{debut} \xrightarrow{w_1} q_i \xrightarrow{a} q_j$. On utilise alors l'hypothèse de récurrence sur w_1 , puis on vérifie que toutes les transitions suivent la loi : $q_j \xrightarrow{0} q_i$ ssi $j * 2 = i$ et $q_j \xrightarrow{1} q_i$ ssi $j * 2 + 1 = i$ (modulo 6)

3. Si w est un mot, on note $\text{mir}(w)$ le mot miroir de w : si $w = a_1.a_2 \dots a_n$ alors $\text{mir}(w) = a_n \dots a_2.a_1$.

Et on définit le langage miroir comme : $\text{mir}(L) = \{\text{mir}(w), w \in L\}$.

Donner un automate pour le langage $\text{mir}(L_0)$.

Correction : Il suffit d'inverser les transitions, de prendre q_0 pour état initial et q_{debut} pour état final. On obtient un automate non-déterministe

En effet, si $w \in L$, alors il existe un chemin allant de q_{debut} à q_0 tel que $q^{i-1} \xrightarrow{a_i} q^i$. Alors en inversant les transitions on a un chemin qui reconnaît $\text{mir}(w)$ dans l'automate B . De même, si on a $w' \in L(B)$, alors il existe un calcul qui, transposé dans A dit que $w' = \text{mir}(w) \in L(A)$. Donc $\exists w' \in L(A), w = \text{mir}(w')$

Exercice 2.10

Donner un algorithme pour calculer l'ensemble des états accessibles d'un automate.

Exercice 2.11

En pratique, le calcul de la partition compatible la plus grossière possible se fait à l'aide d'un tableau ayant autant de lignes et de colonnes que d'états. Il est facile, à l'aide d'un tel tableau, de marquer les états incompatibles entre eux. À l'initialisation, les états non-acceptants seront marqués comme incompatibles avec les états acceptants. Puis on marque récursivement tout couple d'états (p, q) tels que le couple $(transition(p, a), transition(q, a))$ soit marqué pour une certaine lettre $a \in V_t$. Le marquage une fois terminé, la ligne p du tableau de marquage indique les états qui sont dans la même classe de la partition que p : ce sont les états q tels que les couples (p, q) ne sont pas marqués. Formaliser cet algorithme.

Correction :

```

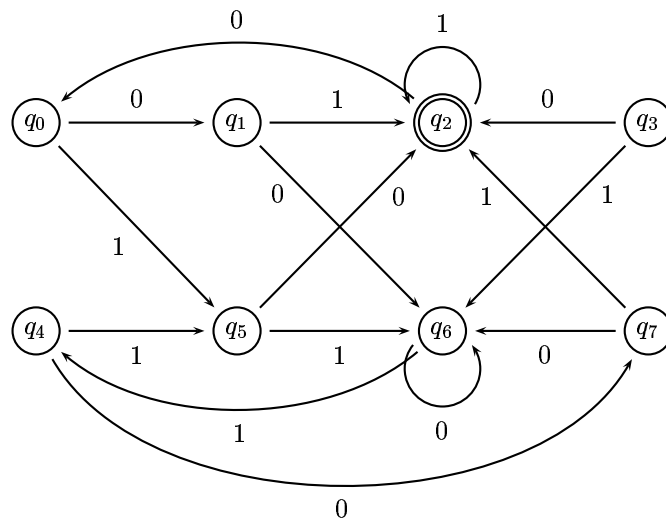
pour chaque  $q \in Q$  faire
  pour chaque  $q \in Q$  faire
    si  $acceptant(p)$  et non  $acceptant(q)$ 
      marquer( $p, q$ )
    fin si
  fin pour
fin pour

répéter
   $stable = true$ 
  pour chaque  $q \in Q$  faire
    pour chaque  $p \in Q$  faire
      pour chaque  $a \in V_t$  faire
        si  $marqué(p, q)$  et non  $marqué(transition(q, a), transition(p, a))$ 
          marquer( $p, q$ )
           $stable = false$ 
        fin si
      fin pour
    fin pour
  fin pour
jusqu'à  $stable$ 

```

Exercice 2.12

Donner l'automate minimal équivalent à l'automate suivant:



Chapitre 3

Expressions rationnelles et analyse lexicale

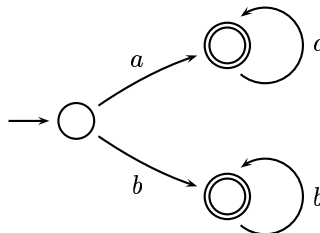
Ce chapitre a pour objectif de décrire la première phase de l'interprétation d'un langage : l'analyse lexicale. Pour ce faire, la classe des langages réguliers étudiée au chapitre précédent va nous servir à définir et à reconnaître les différentes unités lexicales d'un programme, car celles-ci constituent justement des langages réguliers. La définition de ces unités lexicales va nécessiter l'introduction d'un nouveau formalisme de description des langages réguliers : les expressions rationnelles.

Les grammaires régulières ne permettant pas une écriture simple ni concise des langages réguliers (par exemple pour les constantes réelles), nous allons donc en introduire une nouvelle. L'idée que l'on va exploiter pour décrire les langages réguliers de manière simple est d'utiliser les opérations de base des langages : concaténation, itération, union, etc. Nous allons donc étudier les propriétés de clôture des langages réguliers : étant donnés deux automates A_1 et A_2 reconnaissant respectivement les langages L_1 et L_2 , comment peut-on construire un automate reconnaissant $L_1 \times L_2$, L_1^* , $L_1 \cup L_2$, etc ?

Conformément à notre habitude, nous allons d'abord essayer de comprendre ces mécanismes à l'aide d'un exemple simple : supposons $L_1 = \{a\}^+$ et $L_2 = \{b\}^+$ reconnus par les automates suivants :

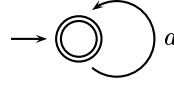


Pour construire l'automate reconnaissant $L_1 \cup L_2$, il suffit alors de réunir les deux états initiaux en un seul ce qui donne l'automate

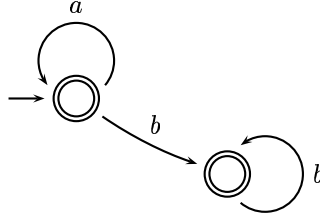


Malheureusement, cette idée de réunir en un seul les deux états initiaux ne fonctionne pas en général : si l'on remplace dans la construction précédente le langage L_1 en le langage

$\{a\}^*$ reconnu par l'automate

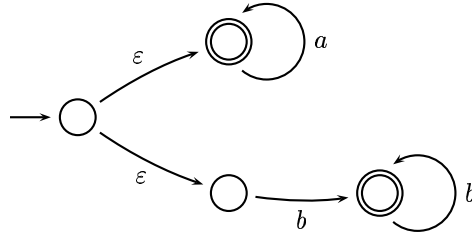


l'automate obtenu par réunion des états initiaux est alors



qui accepte le mot ab alors que celui-ci n'est pas dans la réunion des deux langages. La construction est donc erronée.

Ce phénomène se produit en fait dès l'un des automates possède un circuit passant par son état initial. Pour résoudre ce problème, il nous faut des automates munis d'une opération de choix, permettant de passer d'un état à d'autres états *sans lire une lettre du mot à reconnaître*. Nous allons donc tout d'abord généraliser un peu plus notre définition d'automate en introduisant les automates avec *transitions vides*. La construction correcte de l'union des deux langages précédents sera alors



3.1 Automates non-déterministes avec transitions vides

Un automate avec transitions vides est un automate où certaines transitions peuvent être étiquetées par ϵ au lieu d'une lettre.

Définition 3.1 Un automate non-déterministe A avec transitions vides est un quintuplet (V_t, Q, q_0, F, T) où

1. V_t est le vocabulaire de A ,
2. Q est l'ensemble des états de A ,
3. $q_0 \in Q$ est l'état initial de A ,
4. $F \subseteq Q$ est l'ensemble des états acceptants de A ,
5. T , la fonction de transition de A , est une application de $Q \times (V_t \cup \{\epsilon\}) \rightarrow \mathcal{P}(Q)$. On notera $q \xrightarrow{\alpha} q'$ pour $q' \in T(q, \alpha) = q'$, avec $\alpha \in V_t$ ou $\alpha = \epsilon$.

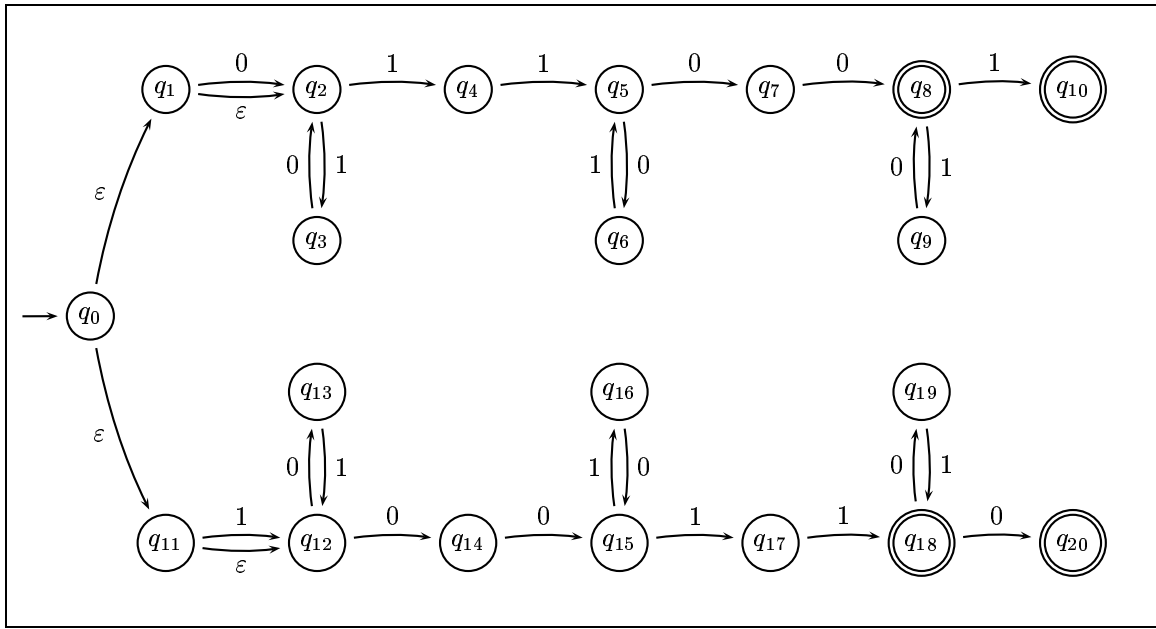


FIG. 3.1 – Un automate non-déterministe avec transitions vides.

L'exécution d'un automate avec transitions vides autorise le passage par un nombre quelconque de transitions vides au cours de l'exécution. Le nombre d'états parcourus n'est bien sûr plus égal au nombre de lettres du mot d'entrée mais peut être plus grand.

Définition 3.2 *Étant donné un automate non-déterministe avec transitions vides $A = (V_t, Q, q_0, F, T)$, on appelle exécution toute suite (éventuellement vide) de transitions $q_0 \xrightarrow{\alpha_1} q_1 \dots q_{n-1} \xrightarrow{\alpha_n} q_n$ avec $\alpha_i \in V_t \cup \{\varepsilon\}$, qui lit le mot $w = \alpha_1 \dots \alpha_n$ (donc de longueur inférieure ou égale à n). Un mot w est reconnu par l'automate A s'il existe une exécution de l'automate A qui se termine en état acceptant en ayant lu le mot w . Comme toujours, on note par $\text{Lang}(A)$ le langage des mots reconnus par l'automate A .*

Nous voulons maintenant montrer que le langage reconnu par un automate avec transitions vides peut également l'être par un automate déterministe. Pour cela, nous allons d'abord éliminer les transitions vides et donc se ramener à un automate non-déterministe, que nous pouvons alors déterminer selon la méthode du chapitre précédent.

Pour éliminer les transitions vides on calcule pour chaque état l'ensemble des états accessibles par le mot vide :

Définition 3.3 *Soit $A = (V_t, Q, q_0, F, T)$ un automate avec transitions vides. La clôture par ε ou en abrégé l' ε -clôture d'un état $q \in Q$ est*

$$\varepsilon\text{-clôture}(q) = \{p \in Q \mid q \xrightarrow{\varepsilon^*} p\}$$

Par exemple, la figure 3.1 représente un automate non-déterministe avec transitions vides reconnaissant le langage sur l'alphabet $\{0, 1\}$ contenant exactement une occurrence de chacun des mots 00 et 11.

On obtient pour les ε -clôtures des états:

état	q_0	q_1	q_2	q_{11}	q_{12}
ε -clôture	$\{q_0, q_1, q_2, q_{11}, q_{12}\}$	$\{q_1, q_2\}$	$\{q_2\}$	$\{q_{11}, q_{12}\}$	$\{q_{12}\}$

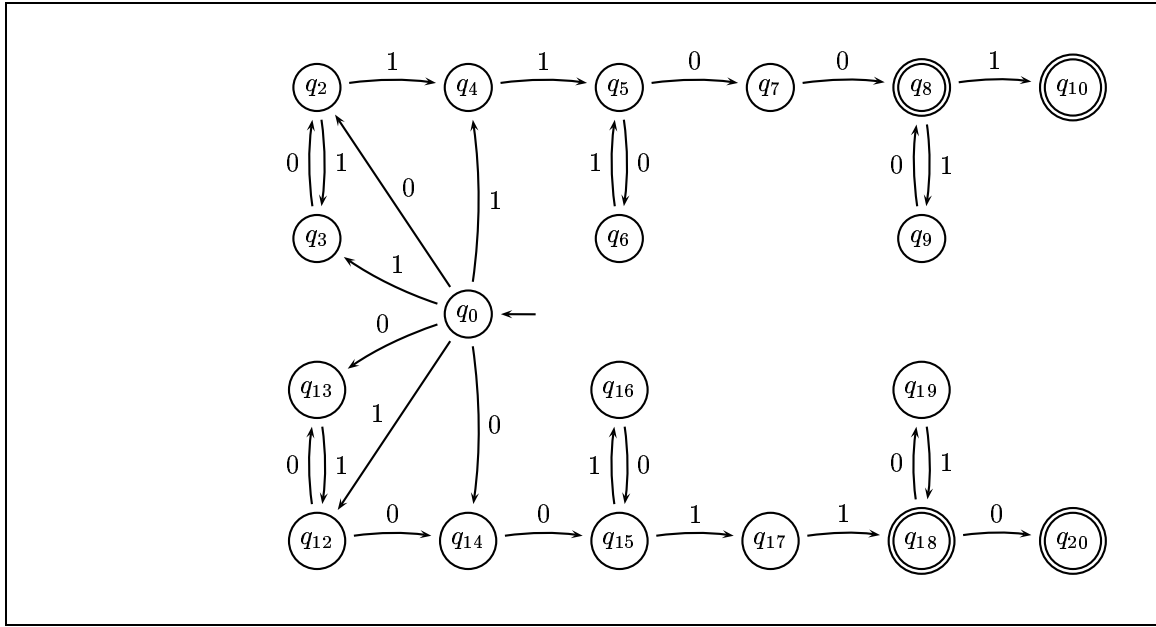


FIG. 3.2 – Automate non-déterministe sans transitions vides.

et $\varepsilon\text{-cl\^oture}(q) = \{q\}$ pour les autres états.

Définition 3.4 Soit $A = (V_t, Q, q_0, F, T)$ un automate non-déterministe avec transitions vides. On définit $E(A)$ comme l'automate $(V_t, Q, q_0, E(F), E(T))$ tel que

$$\begin{aligned} E(T)(q, \alpha) &= \bigcup_{p \in \varepsilon\text{-cl\^oture}(q)} T(p, \alpha) \text{ pour chaque } q \in Q \text{ et } \alpha \in V_t \\ E(F) &= \{q \in Q \mid \varepsilon\text{-cl\^oture}(q) \cap F \neq \emptyset\} \end{aligned}$$

Théorème 3.5 Pour tout automate A non-déterministe avec transitions vides, l'automate $E(A)$ n'a pas de transitions vides et reconnaît le même langage que A .

Ainsi, pour déterminer l'automate de la figure 3.1, on construit d'abord l'automate sans transitions vides mais toujours non-déterministe de la figure 3.2. Nous avons supprimé les états q_1 et q_{11} parce qu'ils ne sont plus accessibles.

L'automate déterminisé correspondant est représenté sur la figure 3.3, où seulement les états accessibles sont montrés. Les états sont étiquetés par l'ensemble des noms des états de l'automate non-déterministe qui le constituent.

Enfin, le minimisé de l'automate de la figure 3.3 est représenté sur la figure 3.4.

3.2 Propriétés de clôture des langages réguliers

Grâce aux automates avec transitions vides, nous sommes maintenant armés pour construire les automates reconnaissant des langages obtenus par concaténation, itération, et union de langages reconnus par des automates.

Pour ces constructions nous allons supposer que tous les automates mis en jeu possèdent un unique état acceptant : on peut toujours se ramener à ce cas au prix de transitions vides additionnelles. Dans les figures qui suivent, un automate arbitraire est figuré par une ellipse

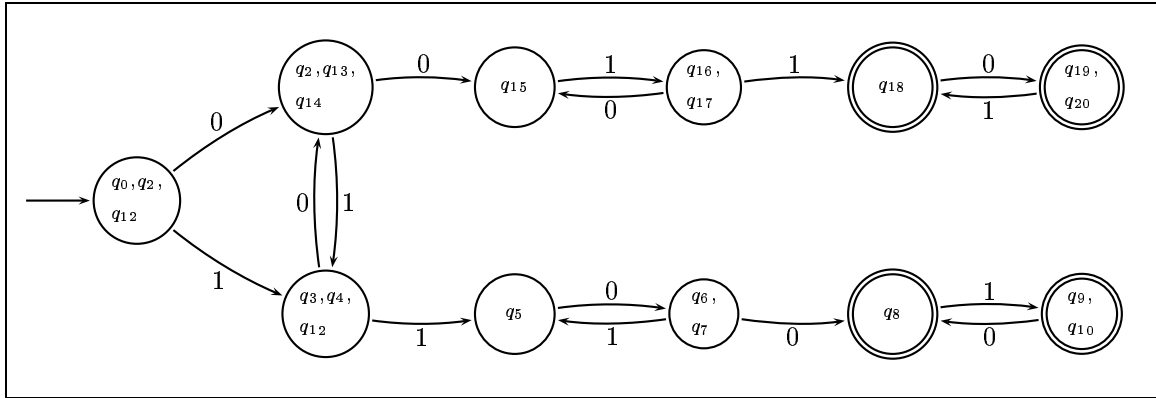


FIG. 3.3 – Automate déterminisé.

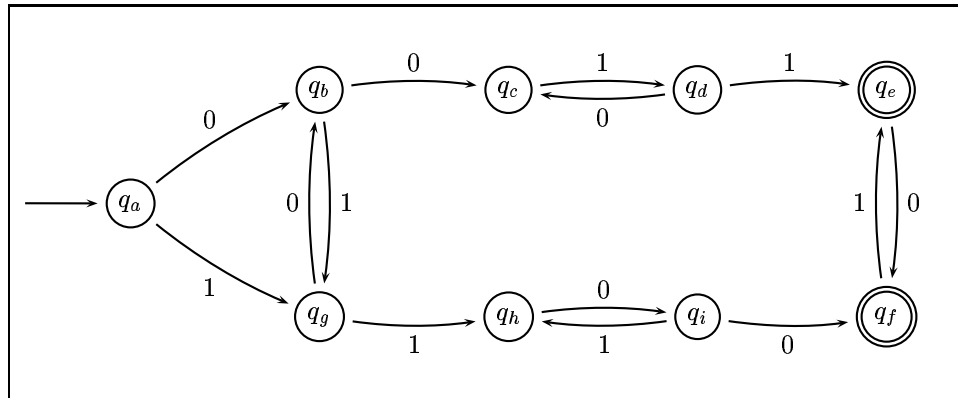
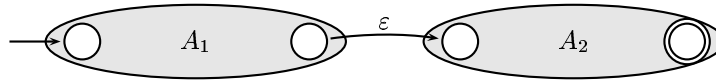


FIG. 3.4 – Automate déterminisé et minimisé.

dans laquelle nous avons fait figurer l'état initial à gauche et l'unique état acceptant à droite. Les constructions conserveront la propriété que l'automate construit possède un unique état acceptant.

3.2.1 Concaténation et Puissance

L'automate reconnaissant la concaténation de deux langages peut être obtenu en ajoutant une transition vide depuis l'état acceptant du premier automate vers l'état initial du deuxième :

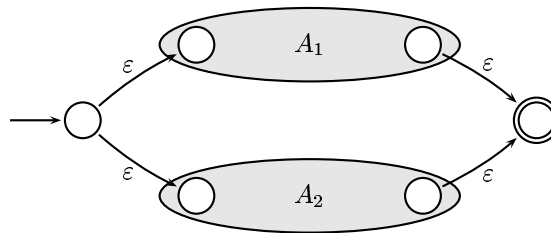


Il faut faire attention à ne pas identifier directement l'état acceptant du premier automate et l'état initial du deuxième, car cela peut permettre de reconnaître plus de mots que nécessaire, dans le cas où il est possible de boucler sur ces états (cf. exercice 3.11).

Il est ensuite facile de construire l'automate reconnaissant la puissance n -ième d'un langage régulier en appliquant $n - 1$ fois la construction précédente.

3.2.2 Union

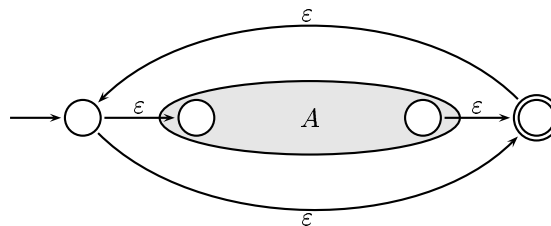
L'automate reconnaissant l'union de deux langages peut être obtenu en « branchant » vers les deux états initiaux. De plus pour conserver la propriété d'unicité de l'état acceptant, nous rajoutons deux transitions vers un nouvel état acceptant :



Toute construction plus simple serait erronée (cf. exercice 3.11).

3.2.3 Itération et itération stricte

L'automate reconnaissant l'itération d'un langage peut être obtenu de la façon suivante :



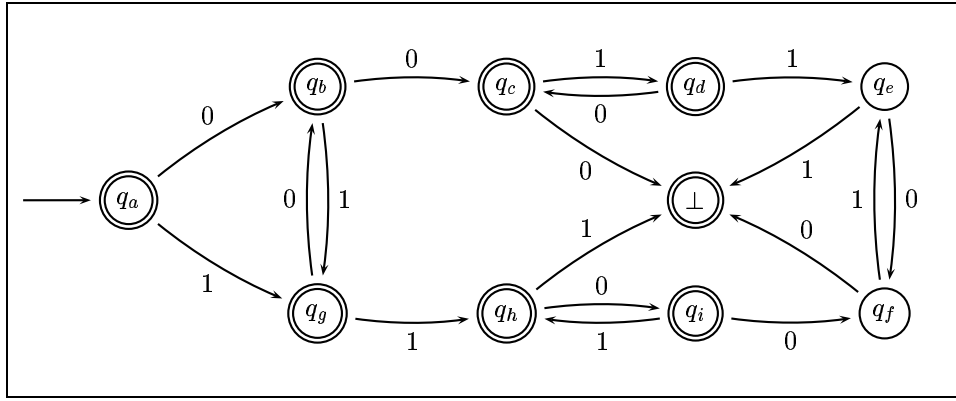


FIG. 3.5 – Complémentation de l'automate de la figure 3.4

Là encore, toute construction plus simple serait erronée (cf. exercice 3.11).

L'automate reconnaissant l'itération stricte d'un langage régulier peut être obtenu par une légère variation de la construction précédente, consistant à supprimer la transition vide allant de l'état initial à l'état final de l'automate construit.

3.2.4 Complémentation et intersection

La complémentation obéit à un principe de construction différent en ce qu'il ne fait pas intervenir de transitions vides. Cette construction suppose donné un automate déterministe complet :

Théorème 3.6 *Soit $A = (V_t, Q, q_0, F, \delta)$ un automate déterministe et complet. Le langage complémentaire $\mathcal{L}ang(A)$ de $\mathcal{L}ang(A)$, est reconnaissable et reconnu par l'automate $A = (V_T, Q, q_0, Q - F, \delta)$.*

En effet, un mot de V_t^* est dans le complémentaire de $\mathcal{L}ang(A)$ si et seulement si l'exécution correspondante n'aboutit pas en un état acceptant de F . La complémentation est donc une opération particulièrement simple, et l'automate obtenu est lui-même déterministe. Par contre, pour calculer l'automate qui reconnaît le complémentaire d'un langage reconnu par un automate quelconque, il faut d'abord déterminer et compléter cet automate.

Il faut bien faire attention que l'opération de complémentation n'est correcte que si l'automate est complet. Par exemple, le complémentaire du langage reconnu par l'automate de la figure 3.4 est l'automate de la figure 3.5, où nous avons rajouté explicitement un état poubelle $\perp$ qui est devenu acceptant par complémentation. Sans cet ajout, l'automate obtenu ne reconnaît pas le complémentaire du langage, contenant en particulier le mot 000.

La clôture par intersection peut maintenant être traitée sans effort par combinaison des clôtures par union et complémentaire (cf. exercice 3.12).

3.3 Expressions rationnelles

En pratique, il n'est pas toujours commode de décrire un langage par un automate, déterministe ou pas. C'est en particulier le cas des automates que l'on manipule en analyse lexicale, qui doivent reconnaître un grand nombre de mots de petite taille. Ces automates

possèdent un grand nombre d'états et sont fastidieux à construire. Il est beaucoup plus facile de décrire un tel langage en énumérant les mots qu'il contient.

Définition 3.7 *L'ensemble des expressions rationnelles $\mathcal{Rat}$ sur l'alphabet V_t est formé comme suit :*

1. tout mot de V_t^* appartient à $\mathcal{Rat}$;
2. si $e \in \mathcal{Rat}$ et $e' \in \mathcal{Rat}$, alors $e \mid e' \in \mathcal{Rat}$;
3. si $e \in \mathcal{Rat}$ et $e' \in \mathcal{Rat}$, alors $e \times e' \in \mathcal{Rat}$;
4. si $e \in \mathcal{Rat}$, alors $e^* \in \mathcal{Rat}$;
5. si $e \in \mathcal{Rat}$, alors $(e) \in \mathcal{Rat}$.

Les expressions rationnelles permettent de dénoter de manière succincte des langages pas trop complexes : il s'agit donc d'un langage permettant de dénoter des langages, de même que le formalisme des grammaires est un langage permettant d'engendrer des langages.

Définition 3.8 *Le langage dénoté par une expression rationnelle e , noté $\mathcal{Lang}(e)$, est défini par :*

1. le mot $w \in V_t^*$ dénote le langage à un mot $\{w\}$;
2. l'expression $e \mid e'$ dénote le langage $\mathcal{Lang}(e) \cup \mathcal{Lang}(e')$;
3. l'expression $e \times e'$ dénote le langage $\mathcal{Lang}(e) \times \mathcal{Lang}(e')$;
4. l'expression e^* dénote le langage $(\mathcal{Lang}(e))^*$;
5. l'expression (e) dénote le même langage que e .

Un langage est dit rationnel s'il peut être dénoté par une expression rationnelle.

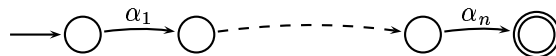
Par exemple, le langage des entiers naturels en représentation binaire est dénoté par l'expression rationnelle $0 + 1 \times (0 + 1)^*$. Se pose à nouveau un problème d'ambiguïté de cette notation, très semblable au problème que nous avons rencontré avec les expressions arithmétiques. Le parenthésage attendu est $0 + (1 \times ((0 + 1)^*))$. Il résulte là aussi de règles de priorité, dans l'ordre décroissant : $*$, $\times$, $+$, et de règles d'association à gauche pour $\times$ et $+$.

Nous avons vu que les langages réguliers et les langages reconnaissables étaient identiques. Il nous reste à les comparer aux langages rationnels. Le résultat fondamental de ce chapitre est le suivant :

Théorème 3.9 *Soit L un langage. Les trois propriétés suivantes sont équivalentes :*

- (i) L est régulier ;
- (ii) L est reconnaissable ;
- (iii) L est rationnel ou bien vide.

La construction d'un automate reconnaissant le langage dénoté par une expression rationnelle se fait grâce aux propriétés de clôtures vues précédemment. Il nous faut tout de même amorcer la construction, en donnant l'automate associé à un mot $\alpha_1 \cdots \alpha_n$:



Le passage d'un automate à une expression rationnelle dénotant le même langage est beaucoup plus complexe et ne sera pas abordé dans ce cours.

3.4 Analyse lexicale avec l'outil LEX

L'analyse lexicale aura pour but de préparer l'étape suivante, l'analyse syntaxique, en engendrant une séquence de *tokens* qui sera ultérieurement analysée par l'analyseur syntaxique. Un analyseur lexical est donc une sorte de gros automate dont les états acceptants reconnaissent les différents tokens (c'est-à-dire les différentes sortes d'unités lexicales) du langage.

Les différents tokens une fois spécifiés sous forme d'expressions rationnelles, un automate déterministe minimal peut être construit pour chacun d'eux. L'outil LEX est un outil disponible dans différents langages de programmation, qui automatise cette transformation.

3.4.1 Notion de token

Lors de l'analyse syntaxique d'un langage, nous avons vu qu'il était pratique de définir la grammaire du langage non pas directement sur le vocabulaire du langage, mais plutôt en supposant déjà définies des grammaires simples (en l'occurrence des grammaires régulières) pour les composants simples du langage appelés *unités lexicales* : les constantes numériques, les identificateurs, les mots-clés, etc. Chaque non-terminal source de l'une de ces grammaires régulières est appelé un *token*. Un token représente donc le langage régulier associé à une unité lexicale donnée.

3.4.2 Fichier de description LEX

En pratique, il est agréable d'utiliser des expressions rationnelles auxiliaires, et la définition de langages rationnels se présente donc sous la forme suivante dans le langage LEX :

1. Une suite de définitions de la forme $D_i = R_i$, où D_i est un nom, et R_i est une expression régulière sur l'alphabet $V \cup \{D_1, \dots, D_{i-1}\}$.
2. Une suite de règles de la forme **Si** P_i **Alors** A_i , où P_i est une expression rationnelle sur l'alphabet $V \cup \{D_i\}_i$ appelée *filtre*, et A_i est un programme caractéristique du type d'unité lexicale reconnu par le filtre P_i .

Par exemple, pour le langage des expressions arithmétiques, on aura :

```
lettre = a | b | c | d | e | f | g | h | i | j | k | l | m | n | o | p | q | r | s | t | u | v | w | x | y | z
chiffre = 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9
```

```
Si chiffre(chiffre)* Alors %traitement des constantes
Si lettre(lettre)* Alors %traitement des identificateurs
Si + Alors %traitement de l'addition
Si × Alors %traitement de la multiplication
```

Les traitements peuvent être en général n'importe quelle expression ou instruction du langage de programmation auquel l'outil LEX est associé. Nous l'utiliserons uniquement pour retourner des *tokens* (voir paragraphe suivant).

Les outils LEX autorisent en général d'autres constructions pour les expressions rationnelles, comme les intervalles de lettres, qui permettent d'abrégier l'écriture. L'outil LEX le plus utilisé est l'original et produit du langage C ; mais il existe également pour d'autres langages, avec quelques variantes. Noter qu'en CAML, l'outil CAMLLEX ne permet pas les définitions.

3.4.3 Description des tokens

LEX va nous servir à reconnaître les différentes unités lexicales sous forme de définitions auxiliaires et de règles. Nous pourrions par exemple avoir un fichier LEX de la forme :

```
lettre = a | b | c | d | e | f | g | h | i | j | k | l | m | n | o | p | q | r | s | t | u | v | w | x | y | z
chiffre = 0 | 1 | 2 | 3 | 4 | 5 | 6 | 7 | 8 | 9
```

Si *if* Alors If

Si *then* Alors Then

Si *else* Alors Else

Si chiffre(chiffre)* Alors Id

Si lettre(lettre)* Alors Cte

Notons que puisque toutes les constantes sont remplacées par le token Cte, on a perdu de l'information. Pour la suite de l'interprétation du langage considéré, il faut garder un lien entre le token et la valeur de la constante correspondante. Le même problème se pose pour les identificateurs.

Il y a plusieurs façons de conserver ce lien, suivant le type de langage utilisé. En programmation fonctionnelle, on peut considérer le token comme étant une fonction dont l'argument est l'unité lexicale qui l'a produit : c'est ce qui est fait dans CAMLLEX. En programmation impérative, il est d'usage de conserver ce lien dans une table, la *table des symboles*. Le token est alors une paire, formée du token lui-même d'une part, et d'un pointeur sur la table des symboles d'autre part. La table des symboles pourra bien sûr contenir un grand nombre d'informations pour chaque occurrence différente d'un même token.

Décider d'une solution ou d'une autre est une question d'implémentation. En ce qui nous concerne, nous supposons que l'unité lexicale représentée par le token t est obtenue par un appel à une fonction val , donc est donnée par l'expression $val(t)$.

3.4.4 Règles de priorité

Étant donné un programme à analyser, c'est-à-dire un mot w sur l'alphabet V du langage de programmation considéré, il y aura en général de multiples façons de découper w en unités lexicales. Un principe de sélection est donc ajouté qui garantit l'unicité du découpage :

- (i) Le filtre choisi est celui qui reconnaît le plus grand préfixe du mot w .
- (ii) Au cas où plusieurs filtres $P_{i_1}, \dots, P_{i_q}$, avec $i_1 < \dots < i_q$, reconnaissent le même (plus grand) préfixe u , le filtre de plus petit indice P_{i_1} est choisi.

Soit par exemple un fragment de langage contenant deux types de tokens, le mot clé « end », et les identificateurs. On aura :

```
lettre = a | b | c | d | e | f | g | h | i | j | k | l | m | n | o | p | q | r | s | t | u | v | w | x | y | z
Si end Alors End
Si lettre(lettre)* Alors Id
```

Les automates correspondants à ces filtres sont représentés sur la figure 3.6 sous la forme d'un automate avec transitions vides dont la partie supérieure traite le mot clé « end » et la partie inférieure traite les identificateurs. À chaque état acceptant peut être associé le token correspondant : End pour q_4 et Id pour q_6 .

Soit « endormi » le mot à reconnaître. Les deux filtres reconnaissent le même préfixe, mais le deuxième reconnaît aussi un préfixe plus long, le second filtre est donc choisi et l'état final de l'exécution de l'automate doit donc être l'état q_6 . Soit maintenant « end » le mot à

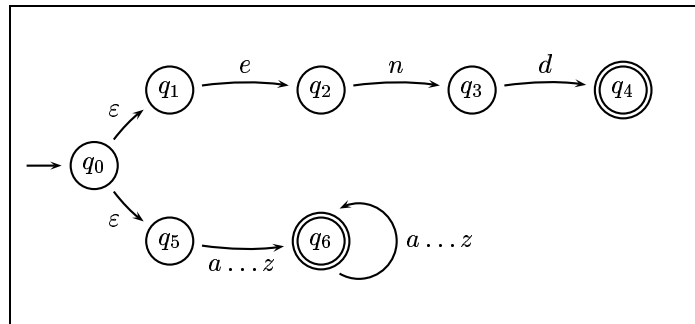


FIG. 3.6 – Automate d'analyse lexicale

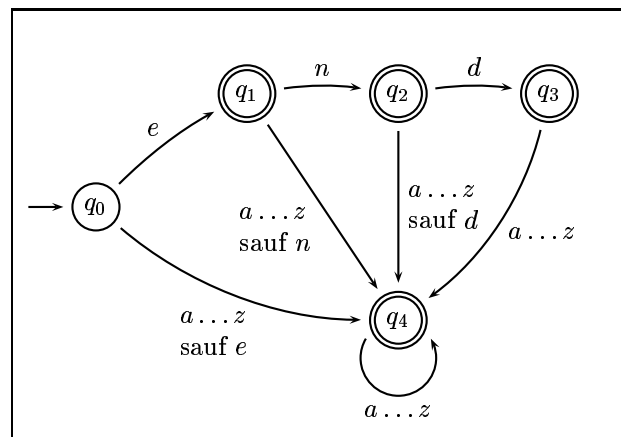


FIG. 3.7 – Automate déterminisé d'analyse lexicale

reconnaître. Les deux filtres reconnaissent le même préfixe de longueur maximale, « end ». Le premier filtre est donc choisi, et l'état final doit donc être l'état q_4 .

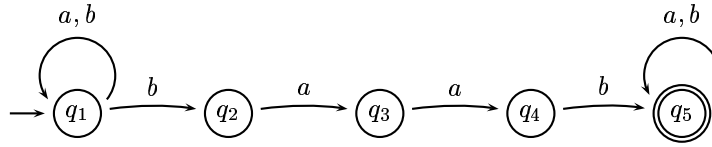
Cet automate est déterminisé (et se trouve être minimal) sur la figure 3.7. L'état acceptant q_3 devrait normalement correspondre aux deux états acceptants de la figure 3.6. Le choix du premier filtre, c'est-à-dire du mot clé, fait que cet état acceptant est en fait associé au mot clé uniquement. Les tokens correspondants aux états acceptants sont donc End pour q_3 et Id pour q_1 , q_2 et q_4 .

En pratique, ces règles de priorité imposent que dans LEX, les filtres correspondants aux différents mots clés devront apparaître avant le filtre pour les identificateurs. De plus, on remarque qu'à chaque état acceptant est associé un unique token, décrit par un filtre P_i , et le code à exécuter lorsque cet état acceptant est atteint est donc le code A_i .

3.5 Exercices

Exercice 3.1

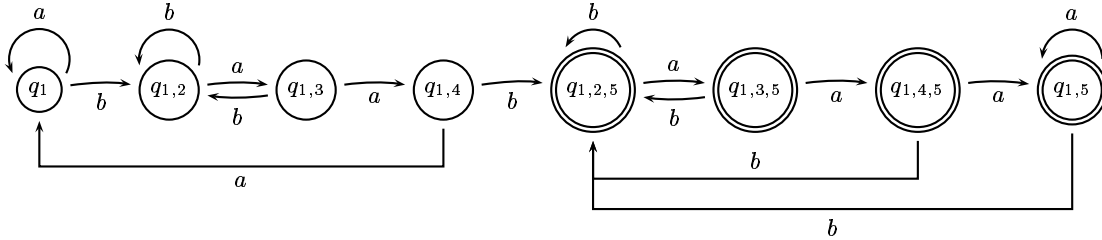
Soit l'automate :



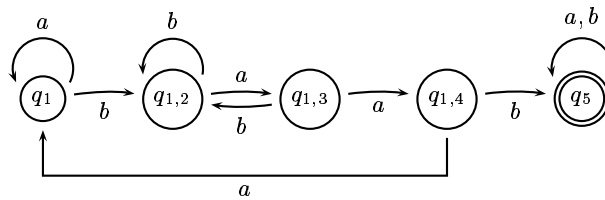
Quel est langage accepté par cet automate? Le déterminer et minimiser.

Correction : Bien sûr, l'automate accepte le langage des mots qui contiennent baab comme sous-mot.

L'automate déterminisé correspondant est :

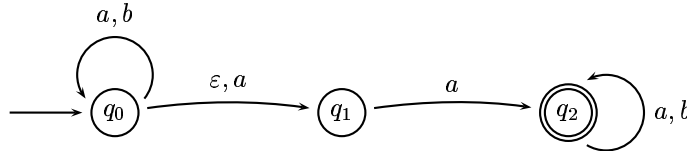


L'automate minimal est :



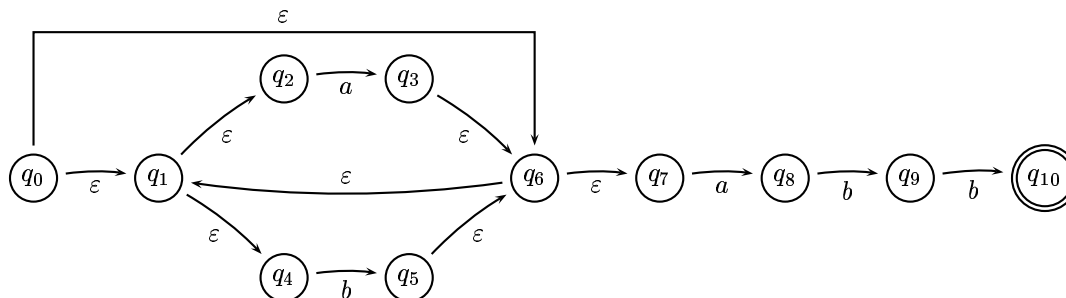
Exercice 3.2

Construire l'automate déterministe équivalent à l'automate suivant :



Exercice 3.3

Soit l'automate :



1. supprimer les transitions vides

Correction :

On supprime d'abord les ε -transitions, pour cela on définit les clotures :

$$\text{cloture}(q_0) = \{q_0, q_1, q_2, q_4, q_6, q_7\} = Q_0$$

$$\text{cloture}(q_1) = \{q_1, q_2, q_4\} = Q_1$$

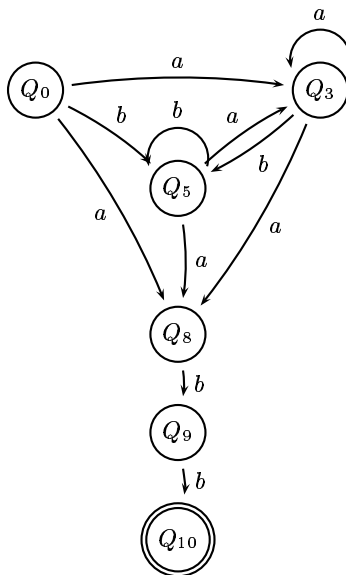
$$\text{cloture}(q_3) = \{q_1, q_2, q_3, q_4, q_6, q_7\} = Q_3$$

$$\text{cloture}(q_5) = \{q_1, q_2, q_4, q_5, q_6, q_7\} = Q_5$$

$$\text{cloture}(q_6) = \{q_1, q_2, q_4, q_6, q_7\} = Q_6$$

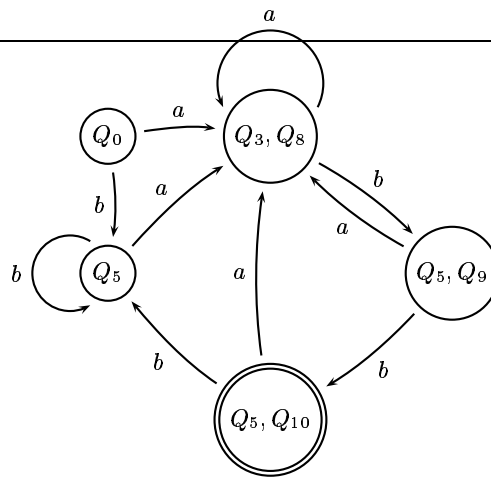
et pour les autres états, on a $\text{cloture}(q_i) = \{q_i\} = Q_i$

L'automate non déterministe sans ε -transition est :



2. Le déterminer

Correction : L'automate déterminisé est :



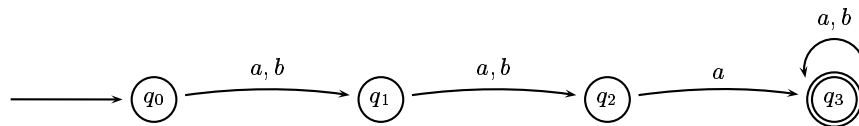
3. Est-il minimal?

Correction : Non, les états Q_0 et Q_5 sont équivalents. L'automate minimal est l'automate sans Q_0 , et avec Q_5 comme état initial

Exercice 3.4

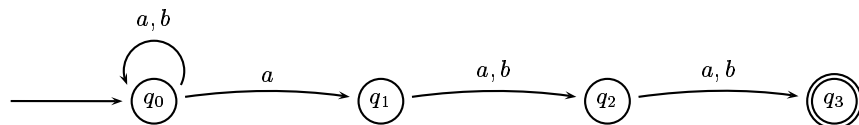
1. Donner un automate déterministe qui accepte le langage sur l'alphabet $\{a, b\}$ dont tous les mots admettent un a comme troisième lettre.

Correction : quatre états suffisent.



2. Donner un automate non-déterministe qui accepte le langage sur l'alphabet $\{a, b\}$ dont tous les mots admettent un a comme troisième lettre à partir de la droite.

Correction : là encore, quatre états suffisent : c'est le miroir du précédent. Mais en faisant le miroir, cela devient non-déterministe.

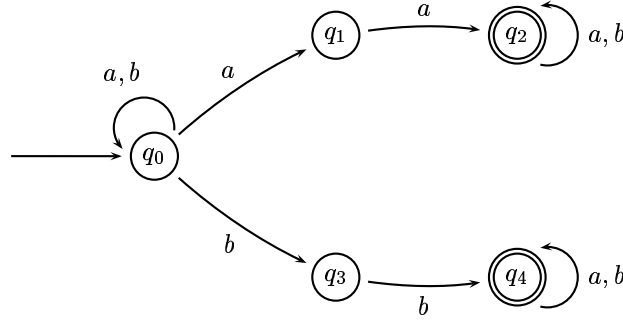


3. Donner un automate déterministe pour le même langage.

Correction : Il a huit états.

Exercice 3.5

Considérons l'automate non-déterministe suivant :

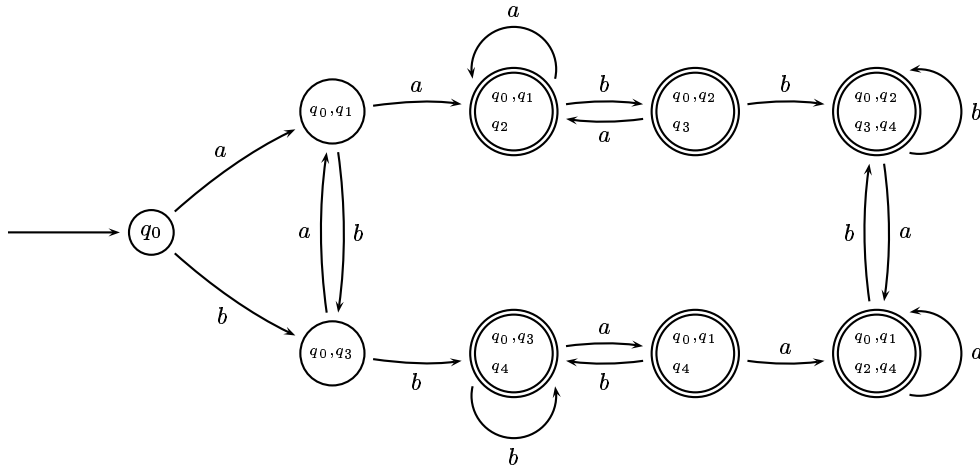


1. Construire un automate déterministe équivalent.

Correction : L'état initial du déterminisé est $\{q_0\}$ On calcule successivement

- (a) $T(\{q_0\}, a) = \{q_0, q_1\}$
- (b) $T(\{q_0\}, b) = \{q_0, q_3\}$
- (c) $T(\{q_0, q_1\}, a) = \{q_0, q_1, q_2\}$
- (d) $T(\{q_0, q_1\}, b) = \{q_0, q_3\}$
- (e) $T(\{q_0, q_1, q_2\}, a) = \{q_0, q_1, q_2\}$
- (f) $T(\{q_0, q_1, q_2\}, b) = \{q_0, q_2, q_3\}$
- (g) etc.

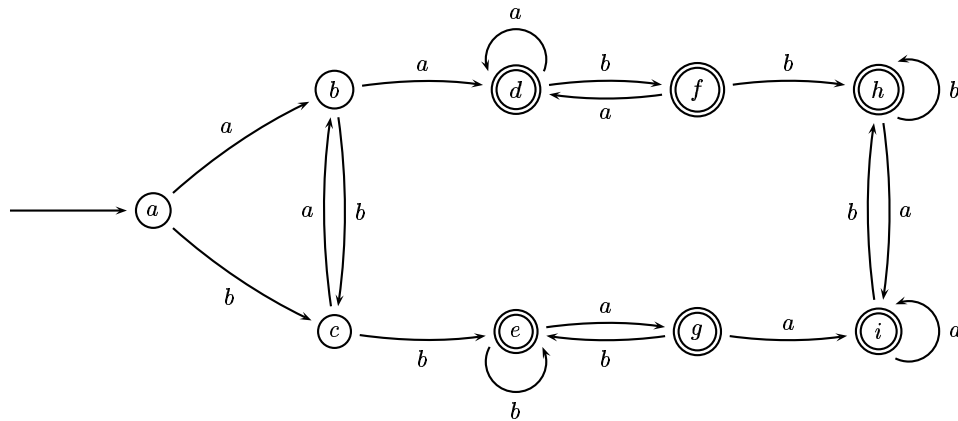
Les états acceptants sont ceux qui contiennent au moins un des états acceptants de l'automate de départ. Ce qui donne



2. Minimiser l'automate obtenu.

Correction :

On renomme les états d'abord pour éviter de se trimballer des noms à rallonge :



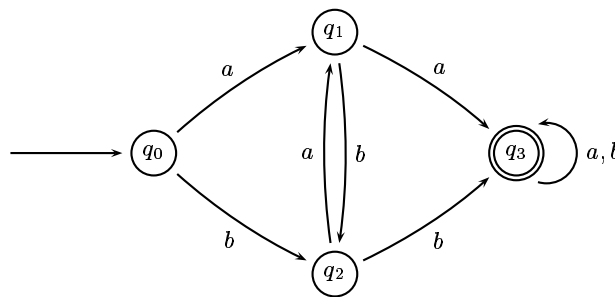
Il faut leur montrer comment on minimise à l'aide d'un tableau : on met une colonne par états, et sur la première ligne on écrit la partition de départ : classe numéro 0 pour les états non acceptants et classe numéro 1 pour les états acceptants. Ensuite, sur chaque ligne successivement :

- on calcule la classe atteinte par lecture de a et de b ;
- on raffine la partition en créant un nouveau numéro de partition à chaque nouveau triplet (numéro de partition, transition par a , transition par b).

Sur l'exercice cela donne :

états	a	b	c	d	e	f	g	h	i
partition	0	0	0	1	1	1	1	1	1
classes de $T(q, a), T(q, b)$	0, 0	1, 0	0, 1	1, 1	1, 1	1, 1	1, 1	1, 1	1, 1
partition	0	1	2	3	3	3	3	3	3
classes de $T(q, a), T(q, b)$	1, 2	3, 2	1, 3	3, 3	3, 3	3, 3	3, 3	3, 3	3, 3
partition	0	1	2	3	3	3	3	3	3

et on s'arrête car c'est stable. D'où l'automate minimal :



Exercice 3.6

Soit L un langage reconnaissable, et L' le langage des mots obtenus à partir des mots de L en supprimant dans les mots de L les lettres de rang impair. Montrer que L' est reconnaissable.

Exercice 3.7

Au café des automates à Paris, quatre consommateurs ayant commandé quatre bières jouent avec un serveur aux yeux bandés et dont les mains sont munies de gants de boxe au jeu

suivant : ils commencent par retourner certains des quatre verres disposés en carré sur la table, puis, à chaque fois que le serveur a lui-même retourné soit un verre, soit deux verres, ils font tourner la table d'un nombre arbitraire de quart de tour. Le serveur étant les yeux bandés et muni de gants de boxe, ne peut savoir si les verres qu'il retourne étaient à l'endroit ou à l'envers : la seule information à sa disposition est que les verres qu'il retourne sont contigus ou opposés. Le jeu s'arrête (avec gain du serveur) s'il arrive à remettre tous les verres dans le même sens, soit à l'endroit, soit à l'envers. Montrer que le serveur possède une stratégie gagnante.

Exercice 3.8

Donner des expressions rationnelles dénotant les langages suivants :

1. les entiers en décimal ;

Correction : si l'on autorise les nombres commençant par 0, cela donne $(0 \mid \dots \mid 9)(0 \mid \dots \mid 9)^*$ ou si l'on est plus exigeant $0 \mid (1 \mid \dots \mid 9)(0 \mid \dots \mid 9)^*$

2. les entiers relatifs en décimal, c'est-à-dire pouvant optionnellement commencer par un signe + ou un signe - ;

Correction : $(+ \mid - \mid \varepsilon)(0 \mid (1 \mid \dots \mid 9)(0 \mid \dots \mid 9)^*)$

3. les identificateurs de PASCAL ;

Correction : $(a \mid \dots \mid z \mid A \mid \dots \mid Z)(a \mid \dots \mid z \mid A \mid \dots \mid Z \mid 0 \mid \dots \mid 9 \mid _)^*$

4. les nombres réels en virgule flottante ;

Correction : $(+ \mid - \mid \varepsilon)(0 \mid (1 \mid \dots \mid 9)(0 \mid \dots \mid 9)^*).(0 \mid \dots \mid 9)^*(0 \mid \dots \mid 9)$

5. les nombres réels en virgule flottante avec exposants.

Correction : $(+ \mid - \mid \varepsilon)(0 \mid (1 \mid \dots \mid 9)(0 \mid \dots \mid 9)^*).(0 \mid \dots \mid 9)^*(0 \mid \dots \mid 9)(\varepsilon \mid (e \mid E)(+ \mid - \mid \varepsilon)(0 \mid (1 \mid \dots \mid 9)(0 \mid \dots \mid 9)^*))$

Exercice 3.9

On dit que deux expressions rationnelles r et s sont *équivalentes* si elles engendrent le même langage. Prouver les équivalences suivantes :

1. $r + s = s + r$
2. $(r + s) + t = r + (s + t)$
3. $(rs)t = r(st)$
4. $r(s + t) = rs + rt$
5. $(r + s)t = rt + st$
6. $\emptyset^* = \varepsilon$
7. $(r^*)^* = r^*$
8. $(\varepsilon + r)^* = r^*$
9. $(r^*s^*)^* = (r + s)^*$

Exercice 3.10

Est-ce que les expressions suivantes sont équivalentes pour n'importe quelles expressions rationnelles r et s :

1. $(rs + r)^*r$ et $r(sr + r)^*$

2. $s(rs + s)^*r$ et $rr^*s(rr^*s)^*$
3. $(r + s)^*$ et $r^* + s^*$

Exercice 3.11

Pour chaque cas de constructions des automates pour la concaténation, l'union et l'itération de langages, donner des exemples d'automates A_1 et A_2 qui la justifie, en ce sens que toute construction plus simple est incorrecte.

Exercice 3.12

1. Exprimer l'intersection de deux langages L_1 et L_2 à l'aide des opérations union et complémentaire de langages.
2. En déduire une méthode de construction de l'automate reconnaissant $L_1 \cap L_2$ en supposant connus les automates reconnaissant L_1 et L_2 .
3. Donner une méthode directe de construction de l'automate reconnaissant $L_1 \cap L_2$ en supposant connus les automates reconnaissant L_1 et L_2 , telle que l'automate obtenu soit déterministe si ces derniers le sont.

Exercice 3.13

Donner un automate reconnaissant le langage des mots formés d'un nombre pair de a et d'un nombre impair de b . En déduire une expression rationnelle pour ce langage.

**Exercice 3.14**

Considérons le langage L des mots *non carrés* sur l'alphabet $\{a, b\}$, c'est-à-dire l'ensemble des mots w tels qu'il n'y ait pas de mot v tel que $w = vv$.

1. Donner une grammaire pour le langage des mots de longueur impaire de L .

Correction : On remarque que tous les mots de longueur impaire sont dans L .

$$A \rightarrow a \mid b \mid aaA \mid abA \mid baA \mid bbA$$

2. Montrer que si $w = v_1av_2bv_3$ avec $|v_2| = |v_1| + |v_3|$ alors $w \in L$.

Correction : Il faut montrer que w n'est pas un carré. Par l'absurde, si $w = vv$ alors $vv = v_1av_2bv_3$ mais comme $|v_2| = |v_1| + |v_3|$, on peut découper v_2 en v_4v_5 avec $|v_4| = |v_3|$ et $|v_5| = |v_1|$, et alors $|v_1av_4| = |v_5bv_3|$ donc $v = v_1av_4 = v_5bv_3$ mais cette dernière égalité est impossible car a et b sont à la même position.

3. En exprimant l'ensemble des mots de la forme ci-dessus (en découpant v_2 de manière astucieuse), donner une grammaire engendrant les mots de longueur paire de L .

Correction : On remarque que l'on peut également découper v_2 en v_4v_5 avec cette fois $|v_4| = |v_1|$ et $|v_5| = |v_3|$. Autrement l'ensemble des mots de la forme ci-dessus est

$$\{v_1av_4v_5bv_3 \mid |v_4| = |v_1| \text{ et } |v_5| = |v_3|\}$$

Ce langage est engendré par la grammaire

$$\begin{aligned} B &\rightarrow B_1B_2 \\ B_1 &\rightarrow a \mid aB_1a \mid aB_1b \mid bB_1a \mid bB_1b \\ B_2 &\rightarrow b \mid aB_2a \mid aB_2b \mid bB_2a \mid bB_2b \end{aligned}$$

On sait que tous les mots de ce langage sont dans L , réciproquement on remarque que les mots de L de longueur paire sont soit de la forme ci-dessus, soit de la forme obtenue en

échangeant a et b . Par conséquent, une grammaire engendrant les mots de longueur paire de L est

$$B \rightarrow B_1 B_2 \mid B_2 B_1$$

4. En déduire une grammaire pour L .

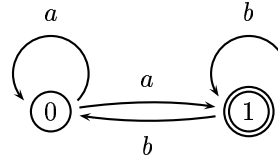
Correction :

$$S \rightarrow A \mid B_1 B_2 \mid B_2 B_1$$

5. Soit L un langage reconnaissable. Montrer que le langage des mots non carrés de L défini par $L' = \{u \mid \text{il n'existe pas de } v \in L \text{ tel que } v = uu\}$ est également reconnaissable.

Exercice 3.15

Soit l'automate $\mathcal{A}$ dessiné ci-dessous:



On demande une expression rationnelle permettant de décrire les langages suivants :

1. Le langage $L_{0,1}^{\{0,1\}}$ des mots reconnus par $\mathcal{A}$ en partant de l'état initial 0, en arrivant en état final 1, et en ne passant jamais par 0 ni 1 entre temps,

Correction : a

2. Le langage $L_{1,0}^{\{0,1\}}$ des mots reconnus par $\mathcal{A}$ en partant de l'état 1 considéré comme état initial, en arrivant en état 0 considéré comme état final, et en ne passant jamais par 0 ni 1 entre temps,

Correction : b

3. Le langage $L_{0,0}^{\{0,1\}}$ des mots reconnus par $\mathcal{A}$ en partant de l'état initial 0, en arrivant en état 0 considéré comme état final, et en ne passant jamais par 0 ni 1 entre temps,

Correction : $\varepsilon + a$

4. Le langage $L_{1,1}^{\{0,1\}}$ des mots reconnus par $\mathcal{A}$ en partant de l'état 1 considéré comme état initial, en arrivant en état final 1, et en ne passant jamais par 0 ni 1 entre temps,

Correction : $\varepsilon + b$

5. Le langage $L_{0,1}^{\{1\}}$ des mots reconnus par $\mathcal{A}$ en partant de l'état initial 0, en arrivant en état 1 considéré comme état final, et en ne passant jamais par 1 entre temps,

Correction : $L_{0,1}^{\{0,1\}} \cup L_{0,0}^{\{0,1\}} L_{0,0}^{\{0,1\}*} L_{0,1}^{\{0,1\}} = a + (\varepsilon + a)(\varepsilon + a)^* a = aa^*$

6. Le langage $L_{1,1}^{\{1\}}$ des mots reconnus par $\mathcal{A}$ en partant de l'état 1 considéré comme état initial, en arrivant en état final 1, et en ne passant jamais par 1 entre temps,

Correction : $L_{1,1}^{\{0,1\}} \cup L_{1,0}^{\{0,1\}} L_{0,0}^{\{0,1\}*} L_{0,1}^{\{0,1\}} = \varepsilon + b + b(\varepsilon + a)^* a = \varepsilon + ba^*$

7. En déduire une expression rationnelle décrivant le langage reconnu par l'automate $\mathcal{A}$

Correction : $L_{0,1}^{\{1\}} L_{1,1}^{\{1\}*} = aa^*(\varepsilon + ba^*)^* = a(a+b)^*$

8. Soit $\mathcal{A}$ un automate fini. On note par $L_{p,q}^I$, avec $I \subseteq Q$, $p \in I$ et $q \in I$, le langage des mots reconnus en allant de l'état p à l'état q dans l'automate $\mathcal{A}$ sans passer par un état intermédiaire appartenant à I . Exprimer $L_{p,q}^I$ en fonction $L_{p,q}^{I \cup \{r\}}$, $L_{p,r}^{I \cup \{r\}}$, $L_{r,r}^{I \cup \{r\}}$ et $L_{r,q}^{I \cup \{r\}}$, où r est un état quelconque n'appartenant pas à I .

Correction : $L_{p,q}^I = L_{p,q}^{I \cup \{r\}} + L_{p,r}^{I \cup \{r\}} L_{r,r}^{I \cup \{r\}}^* L_{r,q}^{I \cup \{r\}}$

9. En déduire un algorithme prenant en entrée un automate $\mathcal{A}$, et rendant comme résultat une expression rationnelle dénotant le langage reconnu par $\mathcal{A}$.

Correction : On utilise la récurrence au-dessus. Il reste le cas $I = Q$:

$$L_{p,q}^Q = \begin{cases} \{a \mid T(p,a) = q\} & p \neq q \\ \{a \mid T(p,a) = q\} \cup \{\varepsilon\} & p = q \end{cases}$$

Ça donne:

$$L = \bigcup_{q \in F} \left(L_{q_0,q}^{\{q\}} L_{q,q}^{\{q\}*} \right)$$

10. Appliquer l'algorithme à l'automate $\mathcal{A}$ de la première question.

Exercice 3.16

1. Donner une expression rationnelle dénotant le langage

$$L_0 = \{a^n b^p \mid \text{où } n \text{ et } p \text{ sont des entiers naturels pairs (y compris 0)}\}$$

Correction : $(aa)^*(bb)^*$

2. Donner un automate déterministe minimal reconnaissant le langage L_0 .

Correction : 4 états, 5 si complet.

3. Donner une expression rationnelle (naïve) dénotant le langage $L = \bigcup_{k \geq 0} L_0^{2k+1}$

Correction : $L = L_0.(L_0.L_0)^*$ donc $(aa)^*(bb)^*((aa)^*(bb)^*(aa)^*(bb)^*)^*$, avec 34 symboles.

4. Donner une expression rationnelle d'au plus 13 symboles qui dénote le langage L (il en existe en fait de moins de 10 symboles). Justifiez votre réponse.

Correction : On commence par constater que $L_0^{2k} \subseteq L_0^{2k+1}$ car $\varepsilon \in L_0$, donc $L = \bigcup_{k \geq 0} L_0^k$, mais comme $\varepsilon \in L_0$, $L = \bigcup_{k \geq 0} L_0^k$, donc L est dénoté par $((aa)^*(bb)^*)^*$ avec 13 symboles. On voit facilement que cette expression est équivalente à $(aa|bb)^*$, avec 8 symboles.

5. En déduire un automate déterministe minimal reconnaissant L .

Correction : Sans pb, avec 3 états, ou 4 si complet.

6. Donner une expression rationnelle dénotant le langage L'_0 des mots obtenus à partir des mots de L_0 en supprimant, pour chacun d'eux, les lettres de rang impair (la première, la troisième, etc.), de telle sorte que le mot $aaaaaabbabbbbbb$ devient $aaabbbb$.

Correction : C'est trivialement a^*b^* .

7. Donner un automate déterministe minimal pour L'_0 .

Correction : 2 états, ou 3 si complet.

8. Donner une expression rationnelle dénotant le langage L_0'' des mots obtenus à partir des mots de L_0 en supprimant, pour chacun d'eux, leur seconde moitié, de telle sorte que le mot $aaaaaabbabbbbbb$ devient le mot $aaaaaab$.

Correction : Si l'on coupe un mot qui a plus de a que de b , on obtient une suite de a , et n'importe quelle suite de a peut être obtenue. Si l'on coupe un mot qui a plus de b que de a , on obtient un nombre pair de a suivi d'un nombre quelconque de b . L_0'' est donc dénoté par $a^*|(aa)^*b^*$.

9. Donner un automate déterministe minimal pour L_0'' .

Correction : 3 états, 4 si complet.

10. Exprimer à l'aide de L et L_0'' le langage L'' des mots obtenus à partir des mots de L en supprimant, pour chacun d'eux, leur seconde moitié, de telle sorte que le mot $aabbbaaabbbaaabbabbbbbb$ devient $aabbbaaabbba$.

Correction : $L = L_0^*$ donc les mots de L'' sont obtenus en découpant un mot de L , de la forme $w_1 \cdots w_n$ à un certain endroit, donc de la forme $w_1 \cdots w_k w$ où w est un mot de L_0 coupé à un endroit a priori quelconque, mais un tel mot est de toute façon dans L_0'' . Donc $L'' = L_0^* \cdot L_0'' = L \cdot L_0''$.

11. En déduire un automate déterministe minimal pour L'' .

Correction : On construit un automate pour L'' en concaténant, avec une transition vide, les automates de L et de L_0'' . La déterminisation et minimisation aboutit à un automate à 3 états, ou 4 si complet. Remarque : une expression rationnelle simple pour L'' est alors $(aa|bb)^*(a|b)$.

Exercice 3.17

1. Donner un automate reconnaissant le langage dénoté par l'expression rationnelle ab^* .

Correction : Faire à l'intuition, pas avec la construction du cours.

2. Donner un automate reconnaissant le langage dénoté par l'expression rationnelle $ab^*(a|b)a$.

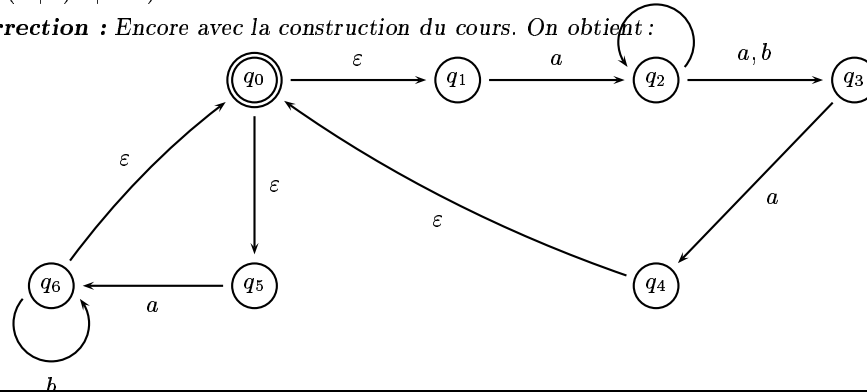
Correction : Idem

3. En déduire un automate reconnaissant le langage dénoté par l'expression rationnelle $ab^*(a|b)a|ab^*$.

Correction : Appliquer la construction du cours. Faire remarquer qu'on ne peut pas se contenter de l'idée naïve consistant à identifier les états initiaux et les états acceptants, qui reconnaîtrait $abab$.

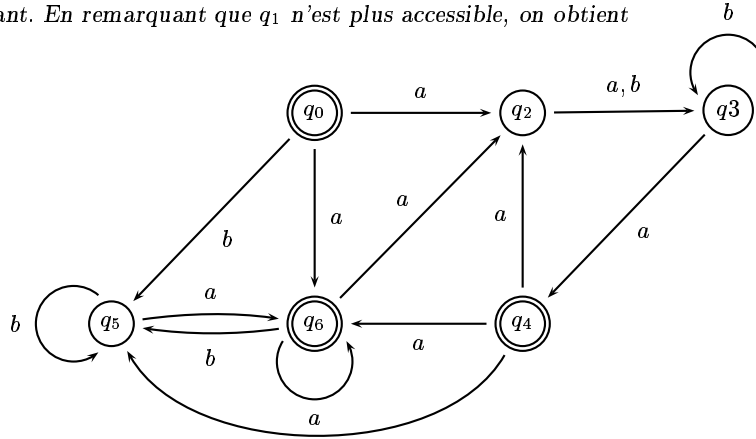
4. En déduire un automate reconnaissant le langage dénoté par l'expression rationnelle $(ab^*(a|b)a|ab^*)^*$.

Correction : Encore avec la construction du cours. On obtient :



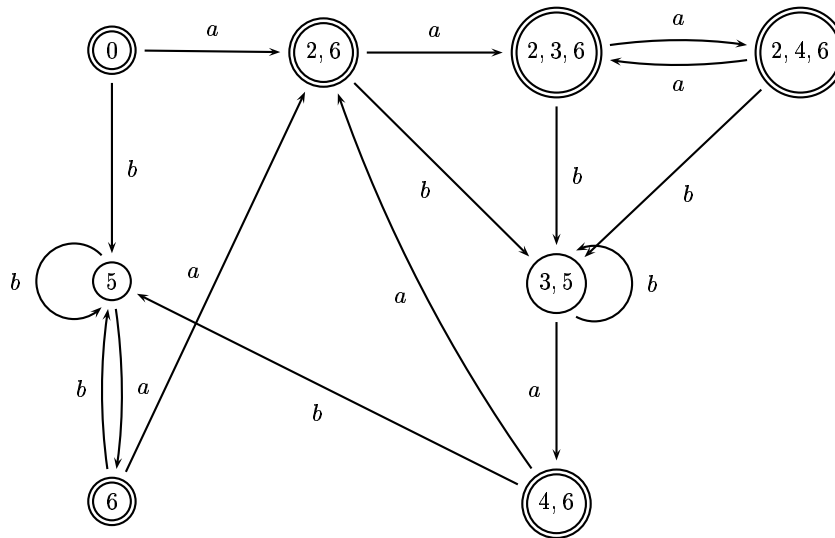
5. Supprimer les transitions vides de l'automate obtenu à la question précédente, puis supprimer les états inaccessibles.

Correction : On commence par calculer les différentes clôtures. $cloture(q_0) = \{q_0, q_1, q_5\}$, $cloture(q_4) = \{q_0, q_1, q_4, q_5\}$, $cloture(q_6) = \{q_0, q_1, q_5, q_6\}$ et pour les autres $cloture(q) = \{q\}$. Les nouveaux états acceptants sont ceux dont la clôture contient un état acceptant. En remarquant que q_1 n'est plus accessible, on obtient



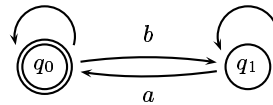
6. Déterminer l'automate obtenu.

Correction :



7. Le minimiser.

Correction : Là, c'est un plaisir puisque l'automate minimal est :



8. Dédire de cet automate une nouvelle expression rationnelle dénotant le langage.

Correction : On fait ça intuitivement puisque la méthode générale est hors programme. Il faut chercher les différents cycles allant de q_0 à q_0 : il y a le cycle passant par a et celui effectuant bb^*a . On obtient donc $(a \mid b.b^*a)^*$.

9. (facultatif) Vérifier par le calcul que ces expressions rationnelles dénotent le même langage. On procédera par inclusion réciproque, et on utilisera la propriété suivante : si $L_1 \subseteq L_2^*$ alors $L_1^* \subseteq L_2^*$.

Correction : On note $L_1 = (L'_1)^* = (a \mid b.b^*a)^*$ et $L_2 = (L'_2)^* = (a(a \mid b)b^*a \mid b^*a)^*$

- $a = b^0a \in L_2$
- $L(bb^*a) \subset L(b^*a)$
Donc $L_1 \subset L_2$
- soit $w \in L(b^*a)$, en regardant sa première lettre, soit $w = a$, soit $w \in L(bb^*a)$ donc $w \in L_2$
- soit $w \in L(a(a \mid b)b^*a)$, alors si w ne contient pas de b il est multiple de a , sinon, il est de la forme $a^1 \text{ ou } ^2bb^{j>0}a$ et se décompose dans L_1

Exercice 3.18

On veut montrer que le langage $L = \{a^n b^n \mid n \in \mathbb{N}\}$ n'est pas régulier.

1. Soit $A = (V_t, Q, q_0, F, T)$ un automate déterministe complet quelconque. Trouver une borne $N \in \mathbb{N}$ associée à A qui assure que pour tout mot $w \in \mathcal{L}ang(A)$ de longueur supérieure ou égale à N , l'exécution de A sur w passe nécessairement au moins deux fois par un certain état.

Correction : Il suffit de prendre $N = |Q|$. En effet, pour tout mot $w = \alpha_1 \cdots \alpha_n$ de longueur $n \geq N$, l'exécution de A sur w a la forme

$$q_0 \xrightarrow{\alpha_1} q_1 \cdots \xrightarrow{\alpha_n} q_n$$

la séquence $q_0, \dots, q_n$ possède $n + 1 > |Q|$ éléments, donc il n'est pas possible que tous les états ne soit parcouru que zéro ou une fois.

2. Raisonnons par l'absurde en supposant que le langage L est reconnaissable par un automate déterministe complet A . Par un raisonnement analogue à la question précédente, trouver un entier N tel que lors de l'exécution de A sur le mot $a^N b^N$, un état est parcouru au moins deux fois lors de la lecture des lettres a .

Correction : Là encore, il suffit de prendre $N = |Q|$. En effet, l'exécution de A sur $a^N b^N$ a la forme

$$q_0 \xrightarrow{a} q_1 \cdots \xrightarrow{a} q_N \xrightarrow{b} q'_1 \cdots \xrightarrow{b} q'_N$$

la séquence $q_0, \dots, q_N$ possède $N + 1 > |Q|$ éléments, donc il n'est pas possible que tous les états ne soit parcouru que zéro ou une fois.

3. En déduire qu'il existe un mot de la forme $a^{N+k} b^N$ avec $k > 0$ reconnu par l'automate, et conclure.

Correction : Dans la séquence précédente, nous savons qu'il existe i et j avec $0 \leq i < j \leq N$ tels que $q_i = q_j$, c'est-à-dire que l'exécution est de la forme

$$q_0 \xrightarrow{a^i} q_i \xrightarrow{a^{j-i}} q_j \xrightarrow{a^{N-j}} q_N \xrightarrow{b^N} q'_N$$

Mézalor nous avons aussi l'exécution suivante qui est possible

$$q_0 \xrightarrow{a^i} q_i \xrightarrow{a^{j-i}} q_j \xrightarrow{a^{j-i}} q_j \xrightarrow{a^{N-j}} q_N \xrightarrow{b^N} q'_N$$

puisque $q_i = q_j$, et donc le mot $a^i a^{j-i} a^{j-i} a^{N-j} b^N = a^{N+j-i} b^N$ est accepté par l'automate. alors qu'il n'est pas dans le langage L , ce qui est une contradiction.

Chapitre 4

Analyse syntaxique

L'objectif de ce chapitre est d'introduire le lecteur aux techniques d'analyse syntaxique et à l'utilisation du logiciel YACC. L'analyse syntaxique est un processus techniquement plus complexe que l'analyse lexicale. Les objets mathématiques sous-jacents, appelés automates à pile, n'ont guère d'applications importantes en dehors de la compilation. Nous ne les étudierons donc pas en détail, nous contentant au contraire de dégager les principes essentiels de l'analyse syntaxique, de façon à comprendre grossièrement le fonctionnement du logiciel YACC.

4.1 Dérivations gauches et droites

Nous allons maintenant décrire le langage des expressions arithmétiques. Nous supposons disposer d'un analyseur lexical reconnaissant le token *Id* décrivant le langage des (noms de) variables, et le token *Cte* décrivant le langage des constantes entières ou réelles (sans les exposants pour simplifier). Ces syntaxes seront supposées être celles des langages informatiques classiques, à l'exception près que les constantes et les variables pourront être de taille arbitraire. Pour simplifier nous considérons également que les seules opérations sont l'addition et la multiplication.

Une grammaire des expressions arithmétiques aura alors pour symboles terminaux $V_t = \{+, \times, (,), Id, Cte\}$ et pour non-terminaux l'unique symbole E qui sera donc la source :

$$E \rightarrow Id \mid Cte \mid E + E \mid E \times E \mid (E)$$

Cette grammaire permet d'engendrer l'expression arithmétique $x + 2.5 \times 4 + (y + z)$, en supposant que les mots 2.5 et 4 font partie du langage engendré par la grammaire des constantes, et que x, y, z font partie du langage engendré par la grammaire des variables. En effet :

$$\begin{aligned} \underline{E} &\rightarrow \underline{E} + E \rightarrow Id + \underline{E} \rightarrow Id + \underline{E} \times E \rightarrow Id + Cte \times \underline{E} \rightarrow Id + Cte \times \underline{E} + E \rightarrow \\ &Id + Cte \times Cte + \underline{E} \rightarrow Id + Cte \times Cte + (\underline{E}) \rightarrow Id + Cte \times Cte + (\underline{E} + E) \rightarrow \\ &Id + Cte \times Cte + (Id + \underline{E}) \rightarrow Id + Cte \times Cte + (Id + Id) \end{aligned}$$

et il suffit de remplacer de manière appropriée les différentes occurrences de *Id* et *Cte* pour obtenir le résultat désiré. La dérivation ci-dessus est appelée *dérivation gauche*, car le non-terminal réécrit à chaque étape est celui qui est situé le plus à gauche dans le mot obtenu à l'étape précédente. La dérivation ci-dessus n'est pas unique, en voici une autre très semblable, qui n'est pas une dérivation gauche, mais est obtenue en permutant les deux dernières étapes

de la dérivation gauche :

$$\begin{aligned} \underline{E} &\rightarrow \underline{E} + E \rightarrow Id + \underline{E} \rightarrow Id + \underline{E} \times E \rightarrow Id + Cte \times \underline{E} \rightarrow Id + Cte \times \underline{E} + E \rightarrow \\ &Id + Cte \times Cte + \underline{E} \rightarrow Id + Cte \times Cte + (\underline{E}) \rightarrow Id + Cte \times Cte + (E + \underline{E}) \rightarrow \\ &Id + Cte \times Cte + (\underline{E} + Id) \rightarrow Id + Cte \times Cte + (Id + Id) \end{aligned}$$

En voici encore une autre, une *dérivation droite* cette fois :

$$\begin{aligned} \underline{E} &\rightarrow E + \underline{E} \rightarrow E + (\underline{E}) \rightarrow E + (E + \underline{E}) \rightarrow E + (\underline{E} + Id) \rightarrow \underline{E} + (Id + Id) \rightarrow \\ &E \times \underline{E} + (Id + Id) \rightarrow \underline{E} \times Cte + (Id + Id) \rightarrow E + \underline{E} \times Cte + (Id + Id) \rightarrow \\ &\underline{E} + Cte \times Cte + (Id + Id) \rightarrow Id + Cte \times Cte + (Id + Id) \end{aligned}$$

Comment être sûr que ces dérivations sont équivalentes, en le sens que les mots engendrés ont la même signification, c'est-à-dire donneront lieu au même calcul ? Il se trouve que le calcul qui est associé à un mot w dépend non seulement du mot lui-même, mais aussi de la dérivation qui a servi à le fabriquer ! Cette dépendance s'exprime par un parenthésage du mot, qui peut être implicite dans le mot w , mais devient explicite dans la dérivation. Par exemple, le parenthésage explicite du mot $x + 2.5 \times 4 + (y + z)$ tel qu'il est défini par la première dérivation est $(x + (2.5 \times (4 + (y + z))))$. La seconde dérivation donne le même parenthésage, mais il n'en est pas de même pour la troisième, qui correspond au parenthésage $((x + 2.5) \times 4) + (y + z)$. En fait le parenthésage attendu n'est ni l'un ni l'autre. Il doit traduire les conventions de *désambiguïsation* habituelles, à savoir que la multiplication a priorité sur l'addition, et que multiplication et addition associent à gauche, c'est-à-dire que $x + y + z$ doit être interprété comme étant $((x + y) + z)$. Cela ne semble pas très important puisque l'addition et la multiplication sont associatives, mais cela le devient lorsque la soustraction et la division font partie du langage (cf. exercice ??). Le bon parenthésage est donc $((x + (2.5 \times 4)) + (y + z))$, et une dérivation de notre exemple conforme à ce parenthésage est la suivante :

$$\begin{aligned} \underline{E} &\rightarrow E + \underline{E} \rightarrow E + (\underline{E}) \rightarrow E + (E + \underline{E}) \rightarrow E + (\underline{E} + Id) \rightarrow \underline{E} + (Id + Id) \rightarrow \\ &E + \underline{E} + (Id + Id) \rightarrow E + E \times \underline{E} + (Id + Id) \rightarrow E + \underline{E} \times Cte + (Id + Id) \rightarrow \\ &\underline{E} + Cte \times Cte + (Id + Id) \rightarrow Id + Cte \times Cte + (Id + Id) \end{aligned}$$

Cette dérivation est elle aussi une dérivation droite. Nous pouvons maintenant en donner une variante (à des permutations près) dans laquelle l'utilisation de la règle $E \rightarrow E \times E$ a lieu lorsqu'aucune autre règle ne peut être utilisée, ce qui est une manière de traduire la convention de désambiguïsation :

$$\begin{aligned} \underline{E} &\rightarrow E + \underline{E} \rightarrow E + (\underline{E}) \rightarrow E + (E + \underline{E}) \rightarrow E + (\underline{E} + Id) \rightarrow \underline{E} + (Id + Id) \rightarrow \\ &\underline{E} + E + (Id + Id) \rightarrow Id + \underline{E} + (Id + Id) \rightarrow Id + E \times \underline{E} + (Id + Id) \rightarrow \\ &Id + \underline{E} \times Cte + (Id + Id) \rightarrow Id + Cte \times Cte + (Id + Id) \end{aligned}$$

Notre mot possède donc plusieurs dérivations droites (et aussi plusieurs dérivations gauches) dans la grammaire, ce qui est la manifestation visible de l'ambiguïté. Nous posons donc la définition suivante.

Définition 4.1 Une grammaire G est ambiguë s'il existe un mot $w \in \mathcal{Lang}(G)$ qui possède plusieurs dérivations droites dans G .

Notre grammaire des expressions arithmétiques est donc ambiguë. Nous allons la transformer de manière à refléter les règles de désambiguïsation, ce qui nécessite l'utilisation de nouveaux non-terminaux F et G . L'idée de la transformation est la suivante : un nouveau symbole non-terminal est utilisé pour chaque niveau de priorité, et une règle récursive

gauche est utilisée pour traduire l'associativité à gauche. Par exemple, la règle $E \rightarrow E + F$ va permettre d'engendrer une suite d'additions parenthésées à gauche. Le non-terminal F servira à engendrer les expressions arithmétiques qui ne sont pas des additions. Il faudra donc rajouter aussi la règle $E \rightarrow F$ de manière à terminer la séquence d'additions. Ainsi, si l'on dispose également de la règle $F \rightarrow a$, on engendrera l'expression $a + a + a$ par la dérivation suivante:

$$\underline{E} \rightarrow \underline{E} + F \rightarrow \underline{E} + F + F \rightarrow \underline{E} + F + F \rightarrow a + \underline{E} + F \rightarrow a + a + \underline{E} \rightarrow a + a + a$$

la règle de multiplication doit être changée conformément à ce principe, et les règles de notre nouvelle grammaire des expressions arithmétiques sont alors les suivantes :

$$E \rightarrow E + F \mid F \qquad F \rightarrow F \times G \mid G \qquad G \rightarrow (E) \mid Cte \mid Id$$

Le même mot est maintenant obtenu par une unique dérivation gauche ou droite, nous donnons ci-dessous la dérivation gauche :

$$\begin{aligned} \underline{E} &\rightarrow \underline{E} + F \rightarrow \underline{E} + F + F \rightarrow \underline{E} + F + F \rightarrow \underline{G} + F + F \rightarrow Id + \underline{E} + F \rightarrow \\ &Id + \underline{E} \times G + F \rightarrow Id + \underline{G} \times G + F \rightarrow Id + Cte \times \underline{G} + F \rightarrow Id + Cte \times Cte + \underline{E} \rightarrow \\ &Id + Cte \times Cte + \underline{G} \rightarrow Id + Cte \times Cte + (\underline{E}) \rightarrow Id + Cte \times Cte + (\underline{E} + F) \rightarrow \\ &Id + Cte \times Cte + (\underline{E} + F) \rightarrow Id + Cte \times Cte + (\underline{G} + F) \rightarrow Id + Cte \times Cte + (Id + \underline{E}) \rightarrow \\ &Id + Cte \times Cte + (Id + \underline{G}) \rightarrow Id + Cte \times Cte + (Id + Id) \end{aligned}$$

La transformation opérée sur la grammaire a éliminé l'ambiguïté au profit d'une part d'une perte de simplicité, car l'écriture de la grammaire est plus complexe, et d'autre part d'une augmentation de la longueur de la dérivation des mots du langage. Nous reviendrons sur cette perte de simplicité au chapitre 5.

4.2 Arbres syntaxiques

Nous avons vu que, dans une même grammaire G , un mot pouvait avoir plusieurs dérivations équivalentes, c'est-à-dire s'obtenant les unes des autres par des permutations adéquates. Nous allons donc maintenant exhiber une structure de donnée, les *arbres syntaxiques*, qui est invariante par ces permutations, de telle sorte qu'un mot de $\mathcal{Lang}(G)$ possède un arbre syntaxique unique lorsque G est non-ambiguë.

Définition 4.2 *Étant donnée une grammaire $G = (V_t, V_n, S, R)$, les arbres de syntaxe de G sont des arbres, les nœuds internes étiquetés par des symboles de V_n , et les feuilles étiquetés par des symboles de V , tels que, si les fils pris de gauche à droite d'un nœud interne étiqueté par le non-terminal N sont étiquetés par les symboles respectifs $\alpha_1, \dots, \alpha_n$, alors $N \rightarrow \alpha_1 \cdots \alpha_n$ est une règle de la grammaire G .*

Un arbre de syntaxe dont la racine est étiquetée par l'axiome de la grammaire et dont le mot des feuilles u appartient à V_t^ est appelé arbre de dérivation de u .*

Tout arbre de dérivation est donc un arbre de syntaxe, mais certains arbres de syntaxe ne sont pas des arbres de dérivation, soit parce que leur racine n'est pas étiquetée par l'axiome de la grammaire, soit parce que certaines de leur feuilles sont étiquetées par des non-terminals.

L'arbre de dérivation définit tout un ensemble de dérivations de son mot des feuilles. Une dérivation est en fait obtenue dès que l'on a fixé un mode de parcours de l'arbre dans lequel un nœud ne peut être visité que si son père l'a déjà été. Ainsi, la figure 4.1 montre deux arbres de dérivation du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire ambiguë des expressions

les arbres de dérivation sont également appelés *arbres syntaxiques*.

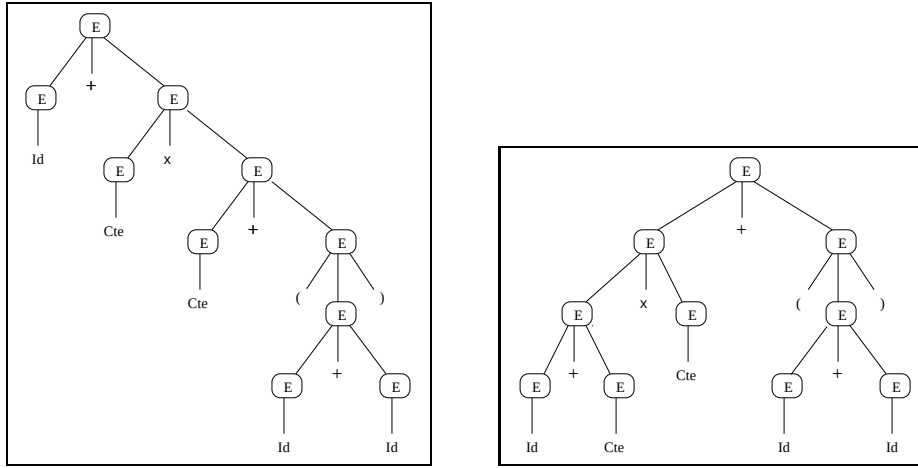


FIG. 4.1 – Deux arbres de dérivation du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire ambiguë des expressions arithmétiques

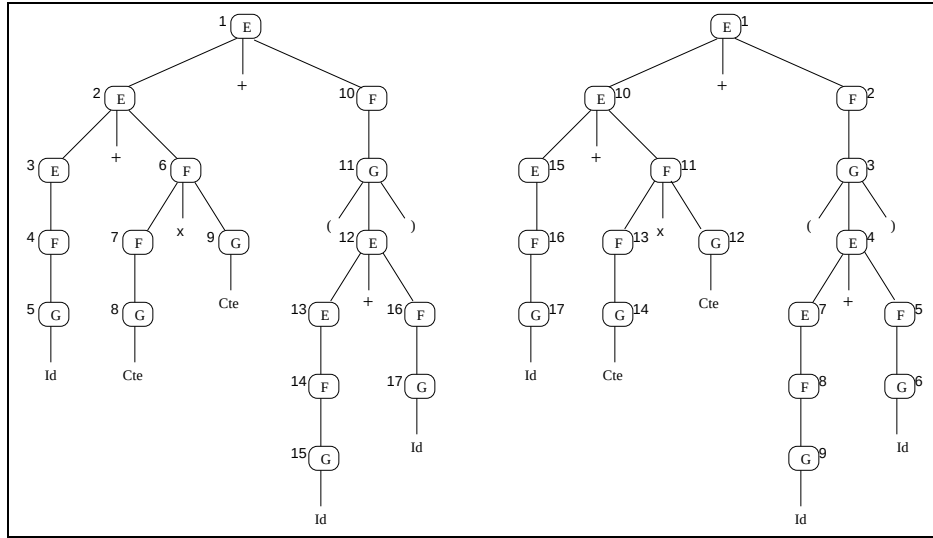


FIG. 4.2 – Parcours gauche et droit de l'arbre de dérivation du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire des expressions arithmétiques désambiguïsée

arithmétiques. Partant de la racine, le parcours gauche consiste à descendre tant que cela est possible en choisissant en permanence la branche la plus à gauche. Lorsque l'on arrive à une feuille, il faut remonter au plus proche ancêtre dont un fils n'a pas encore été visité, puis recommencer la descente à partir du plus à gauche des fils non encore visités. La figure 4.2 indique les parcours gauches et droits de l'arbre de dérivation du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire désambiguïsée, en numérotant les nœuds internes de l'arbre de dérivation dans l'ordre où ils sont visités.

La discussion précédente montre que le mot des feuilles d'un arbre de dérivation de la grammaire G appartient à $\mathcal{Lang}(G)$. Nous allons maintenant montrer la réciproque en construisant l'arbre de syntaxe associé à une dérivation. Notons que cette fois la dérivation est supposée quelconque : son origine est un non-terminal arbitraire, et le mot engendré peut

comporter des non-terminaux. On va donc construire des arbres de syntaxe qui ne seront pas nécessairement des arbres de dérivation. On procède par récurrence sur la longueur des dérivations de la façon suivante :

- à une dérivation de longueur nulle, c'est-à-dire à un symbole $\alpha \in V$, on associe un arbre de syntaxe réduit à une feuille étiquetée par α ;
- à une dérivation de longueur $n + 1$, de la forme $N \rightarrow^* w_1 \underline{M} w_2 \rightarrow w_1 \alpha_1 \dots \alpha_n w_2$, on associe l'arbre de syntaxe obtenu à partir de l'arbre de syntaxe de la dérivation $N \rightarrow^* w_1 M w_2$ en ajoutant, à la feuille étiquetée par M , n feuilles étiquetées respectivement par $\alpha_1, \dots, \alpha_n$.

On en déduit la propriété fondamentale suivante :

Théorème 4.3 *Étant donnée une grammaire $G = (V_t, V_n, S, R)$, un mot $u \in V_t^*$ appartient à $\text{Lang}(G)$ si et seulement s'il possède un arbre de dérivation dans la grammaire G .*

On remarque facilement que des dérivations qui s'échangent par des permutations adéquates ont le même arbre de dérivation. On peut donc reformuler notre définition d'ambiguïté de la manière suivante :

Définition 4.4 *Une grammaire G est ambiguë s'il existe un mot $u \in \text{Lang}(G)$ qui possède deux arbres de dérivation distincts dans la grammaire G .*

La figure 4.1 confirme que notre première grammaire des expressions arithmétiques était ambiguë.

4.3 Applications à l'analyse syntaxique

Il existe en fait deux façons bien différentes de concevoir l'analyse syntaxique d'un mot d'un langage décrit par une grammaire hors contexte. Lors de la lecture du mot (de la gauche vers la droite), l'arbre syntaxique peut être construit de haut en bas, ou de bas en haut, conformément à la figure 4.3. Nous allons maintenant illustrer ces deux méthodes sur des exemples.

4.3.1 Analyse syntaxique descendante

Considérons la grammaire de Dick :

$$S \rightarrow \varepsilon \mid (S)S$$

L'analyse descendante correspond à une dérivation gauche. Par exemple, pour le mot $((()())())$, la dérivation gauche est la suivante :

$$\begin{aligned} S &\rightarrow (S)S \rightarrow ((S)S)S \rightarrow (((S)S)S)S \rightarrow (((()S)S)S)S \rightarrow (((()S)S)S)S \rightarrow (((()S)S)S)S \\ &\rightarrow (((()S)S)S)S \rightarrow (((()S)S)S)S \rightarrow (((()S)S)S)S \rightarrow (((()S)S)S)S \rightarrow (((()S)S)S)S \end{aligned}$$

L'analyse syntaxique descendante va construire la suite des arbres d'analyse correspondant à cette dérivation, représentée à la figure 4.4.

Le choix entre les deux règles applicables peut se faire grâce à la connaissance du prochain caractère devant être engendré, appelé le *caractère d'avance*. Au départ, ce caractère est la parenthèse ouvrante, et la première règle à appliquer doit donc engendrer une parenthèse ouvrante. Il se trouve que la seconde règle engendre une parenthèse ouvrante, et ce sera donc elle que l'on appliquera chaque fois qu'une parenthèse ouvrante devra être engendrée. De

son coté, la première règle fait disparaître le non-terminal S . Or, dans le membre droit de la seconde règle, la première occurrence du non-terminal S est placée devant une parenthèse fermante. L'application de la première règle à une telle occurrence de S engendre donc indirectement une parenthèse fermante. Il en est de même de la seconde occurrence de S (il faut pour s'en rendre compte regarder les arbres syntaxiques) sauf dans le cas où il s'agit de la dernière telle occurrence de S à remplacer avant de terminer (et alors il faut ne rien engendrer du tout, car le mot à lire est entièrement lu). Le choix de l'application d'une règle est donc entièrement déterminé par la connaissance du caractère d'avance. Une telle grammaire est dite LL(1), le premier L étant l'initiale du mot anglais Left, qui indique que la lecture du mot à analyser se fait de gauche à droite, le second L, initiale du même mot anglais, indiquant que l'on construit une dérivation gauche, et le 1 indiquant l'utilisation d'un unique caractère d'avance. Toutes les grammaires ne sont pas LL(1), certaines peuvent être LL(2), sans être LL(1), mais seul le cas LL(1) est utilisé en pratique. Se rendre compte qu'une grammaire est LL(1) demande le calcul des premiers caractères engendrés par l'utilisation de chaque règle. Par exemple, dans la grammaire de Dick, le premier caractère engendré par l'utilisation de la seconde règle est la parenthèse ouvrante. Ces calculs sont simples, mais nous nous contenterons de cette description intuitive. Notons enfin qu'une grammaire LL(1) est nécessairement non-ambiguë, puisque le choix d'une règle lors de l'analyse est entièrement déterminé par la connaissance du caractère d'avance.

4.3.2 Analyse syntaxique ascendante

L'analyse descendante est très simple à mettre en œuvre, puisqu'il s'agit tout simplement de construire une dérivation gauche. Malheureusement, elle échoue dans de très nombreux cas. Elle échoue en particulier chaque fois que la grammaire est « récursive gauche », c'est-à-dire qu'un membre droit de règle commence par le non-terminal qui en est le membre gauche. Considérons par exemple la grammaire désambiguïsée des expressions arithmétiques :

$$\begin{aligned} E &\rightarrow E + T \mid T \\ T &\rightarrow T \times F \mid F \\ F &\rightarrow (E) \mid Id \mid Cte \end{aligned}$$

Les règles $E \rightarrow E + T$ et $T \rightarrow T \times F$ sont récursives gauches. De plus, la technique précédente ne permettrait visiblement pas de choisir entre les règles $E \rightarrow E + T$ et $E \rightarrow T$. L'analyse ascendante que nous allons voir maintenant n'a pas cet inconvénient, mais elle est d'une mise en œuvre plus complexe. Le principe est de construire une dérivation droite à l'envers. La construction d'une dérivation droite à l'endroit nécessiterait en effet une lecture de droite à gauche du mot à analyser. La lecture de gauche à droite est donc responsable de cette construction à l'envers de la dérivation. Ce mode opératoire a pour conséquence que les objets manipulés au cours de l'analyse ne sont pas des arbres de syntaxe, mais des séquences d'arbres de syntaxe appelées *forêts*. Une forêt a pour *racine* la séquence des racines des arbres qui la composent, c'est-à-dire un mot sur l'alphabet $V_t \cup V_n$ de la grammaire. Par exemple, pour le mot $Id \times Id + Cte$, la suite des forêts construites est représentée à la figure 4.5. La suite des racines de ces forêts est la suite de mots $Id, F, T, T \times, T \times Id, T \times F, E, E +, E + Id, E + F, E + T, E$.

Les opérations effectuées au cours de l'analyse ascendante, comme on peut le remarquer sur la figure 4.5, sont de deux sortes :

- (i) la lecture d'un caractère, opération appelée *shift*.
- (ii) l'enracinement d'une règle, opération appelée *reduce*. Cette opération consiste à écrire la forêt courante g sous la forme de la juxtaposition de deux forêts ff' , puis à construire une nouvelle forêt en juxtaposant la forêt f avec l'arbre de syntaxe de racine N , membre gauche d'une règle $N \rightarrow racine(f')$, et dont f' est la suite des sous-arbres.

qui n'apparaît pas dans un membre droit de règle. Des conflits shift-reduce se présentent également, à chaque fois que la racine de la forêt courante se termine par $E + T$ ou $T \times F$. On peut alors enraciner avec les règles $E \rightarrow T$ et $E \rightarrow E + T$ dans le premier cas, et avec les règles $T \rightarrow F$ et $T \rightarrow T \times F$ dans le second. Dans les deux cas, on préférera la règle de plus long membre droit, c'est à dire $E \rightarrow E + T$ ou $T \rightarrow T \times F$. De nouveau, le choix contraire interdirait l'enracinement ultérieur de $+E$ dans le premier cas et de $\times T$ dans le second. Comme il est possible de résoudre tous les conflits, la grammaire des expressions arithmétiques est LR(1). Notons là-encore qu'une grammaire LR(1) est non-ambiguë.

4.4 Génération d'analyseurs syntaxiques avec YACC

On dispose aujourd'hui d'outils qui permettent la génération d'analyseurs syntaxiques à partir de grammaires hors-contextes, dont le plus connu est YACC. Étant donnée une grammaire, YACC essaye de trouver une stratégie pour une analyse ascendante en utilisant la méthode décrite dans la section précédente. La classe de grammaires pour lesquelles YACC est capable d'engendrer un analyseur est une variante de LR(1) appelée LALR(1). En fait, YACC accepte comme données des grammaires qui peuvent être ambiguës (donc ne sont exactement pas LALR), pourvu qu'elles soient accompagnées des *règles de désambiguïsation* adéquates qui précisent les priorités respectives des différents opérateurs. Cette facilité permet de garder des grammaires simples, mais son utilisation nécessite une bonne compréhension du fonctionnement du logiciel. Nous nous interdirons pour cette raison cette possibilité de YACC : les grammaires que nous utiliserons devront toujours être non-ambiguës.

Dans la section précédente, nous avons décrit le principe de l'analyse ascendante comme une construction explicite de l'arbre de dérivation. Pour la réalisation d'un tel algorithme il n'est pas du tout nécessaire de garder toutes les forêts construites : comme toutes les actions de l'analyseur dépendent seulement du prochain caractère lu et des *racines* des forêts construites, YACC va économiser l'utilisation de la mémoire et stocker seulement les racines des forêts. Autrement dit, l'analyseur engendré par YACC par défaut fournit seulement une fonction disant si un mot donné est dans le langage engendré par la grammaire ou pas. Cependant, l'analyse syntaxique n'est en principe qu'un premier pas dans un logiciel plus complexe, et on a donc généralement besoin d'obtenir un résultat plus informatif. Dans ce but, YACC permet d'associer aux règles de la grammaire des expressions permettant de construire le résultat souhaité de l'analyse syntaxique, arbre de syntaxe, résultat d'évaluation, etc. Ces expressions prennent des valeurs dans une domaine défini dans le langage de programmation, et utilisent des opérations prédéfinies (ou définies par le programmeur) dans ce langage. L'exemple ci-dessous décrit une grammaire des expressions arithmétiques, ainsi que l'évaluation des expressions reconnues (notre syntaxe dérive de la syntaxe réelle de YACC, à trouver dans un manuel de référence) :

Non-terminaux : E, T

Axiome : E

Tokens : Cte avec valeur de type entier, $+$, $($, $)$

$$\begin{array}{lll} E & \rightarrow & E + T \quad \{E + T\} \\ & | & T \quad \{T\} \\ T & \rightarrow & (E) \quad \{E\} \\ & | & \text{Cte} \quad \{\text{val}(\text{Cte})\} \end{array}$$

On fournit donc à YACC les informations concernant la grammaire d'une part, c'est-à-dire la liste des non-terminaux, l'axiome, la liste des tokens, celle des règles, et d'autre part un mécanisme d'évaluation donné sous la forme d'expressions accolées aux règles. L'expression $\{E + T\}$ accolée à la règle $E \rightarrow E + T$ spécifie un calcul à effectuer sur l'arbre

de syntaxe. La règle $E \rightarrow E + T$ signifie que l'arbre de syntaxe de racine E est formé d'un arbre de racine E (le E du membre gauche de règle), dont les trois sous-arbres ont pour racines respectives E (le E du membre droit de règle), $+$ et T . Notons par A_1 et A_2 les deux sous-arbres de syntaxe dont les racines E et T sont des non-terminaux. Supposons (c'est notre hypothèse de récurrence) calculées les valeurs v_1 et v_2 des sous-arbres A_1 et A_2 . Le calcul de la valeur de l'arbre tout entier s'exprime comme une fonction, qui ne dépend que de la règle $E \rightarrow E + T$, dont les arguments sont v_1 et v_2 . Dans le cas présent, ce sera la fonction qui calcule la somme $v_1 + v_2$. Quand il y a plusieurs occurrences du même token ou non-terminal dans le membre droit d'une règle de la grammaire, on doit les numérotérer pour les distinguer, comme dans l'exemple suivant, qui calcule le nombre de paires de parenthèses dans un mot de la grammaire de Dick à une parenthèse :

Non-terminaux: S
 Axiome: S
 Tokens: $(,)$

$$\begin{array}{lcl} S & \rightarrow & \varepsilon \quad \{0\} \\ & | & (S)S \quad \{S_1 + S_2 + 1\} \end{array}$$

Ce mécanisme est utilisé de manière systématique pour calculer une représentation compacte de l'arbre de dérivation, comme expliqué dans le chapitre 5.

4.5 Exercices

Exercice 4.1

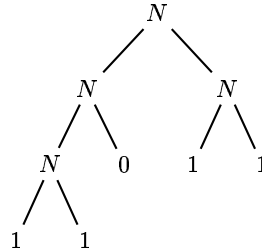
On considère la grammaire suivante :

$$G = (\{0, 1\}, \{N\}, N, \{N \rightarrow 11 \mid 1001 \mid N0 \mid NN\})$$

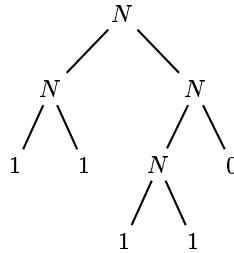
1. Dessiner les arbres de dérivation pour les mots 11011 et 11110.

Correction :

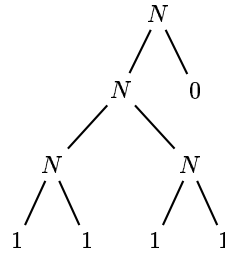
- $\underline{N} \rightarrow \underline{N}N \rightarrow \underline{N}0N \rightarrow 110\underline{N} \rightarrow 11011$.



- $\underline{N} \rightarrow \underline{N}N \rightarrow 11\underline{N} \rightarrow 11\underline{N}0 \rightarrow 11110$



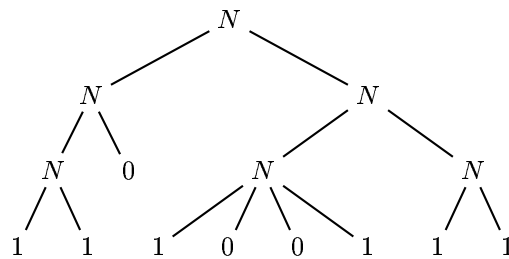
ou $\underline{N} \rightarrow \underline{N}0 \rightarrow \underline{N}N0 \rightarrow 11\underline{N}0 \rightarrow 11110$



2. Est-ce que la grammaire est ambiguë?

Correction : Oui. Voir les dérivations du second mot ci-dessus.

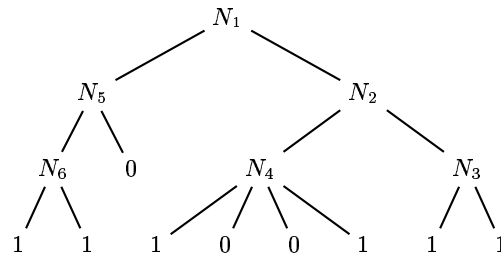
3. On regarde l'arbre syntaxique suivant:



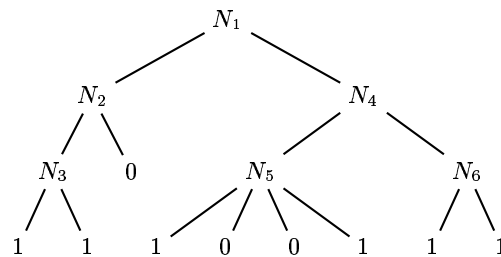
Donner la dérivation gauche et la dérivation droite pour cet arbre.

Correction :

- Doite: $\underline{N} \rightarrow \underline{N}\underline{N} \rightarrow \underline{N}\underline{N}\underline{N} \rightarrow \underline{N}\underline{N}11 \rightarrow \underline{N}100111 \rightarrow \underline{N}0100111 \rightarrow 110100111$.



- Gauche: $\underline{N} \rightarrow \underline{N}\underline{N} \rightarrow \underline{N}0\underline{N} \rightarrow 110\underline{N} \rightarrow 110\underline{N}\underline{N} \rightarrow 1101001\underline{N} \rightarrow 110100111$.



Exercice 4.2

La grammaire G de l'exercice 1.15, est-elle ambiguë?

Correction : Oui, on considère le mot $aabbab$.

Exercice 4.3

Trouver une grammaire pour le langage suivant :

$$\{a^i b^j c^k \mid i, j, k \geq 0, i = j \text{ ou } j = k\}$$

Votre grammaire est-elle ambiguë ? Pourquoi ?

Correction :

$$\begin{aligned} S &\rightarrow XY \mid WZ \\ X &\rightarrow aXb \mid \varepsilon \\ Y &\rightarrow cY \mid \varepsilon \\ W &\rightarrow aW \mid \varepsilon \\ Z &\rightarrow bZc \mid \varepsilon \end{aligned}$$

La grammaire est nécessairement ambiguë car dans le cas où le nombre de a , de b et de c sont égaux, il y a ambiguïté dans la définition : soit le mot est dans le langage parce que $i = j$ soit parce que $j = k$. On peut s'en convaincre avec abc qui possède deux dérivations gauches distinctes :

$$\begin{aligned} S &\rightarrow XY \rightarrow aXbY \rightarrow abY \rightarrow abc \\ S &\rightarrow WZ \rightarrow aZ \rightarrow abZc \rightarrow abc \end{aligned}$$

Exercice 4.4

On considère le langage des expressions arithmétiques avec addition, opérateur infixé noté $+$, soustraction, opérateur infixé noté $-$, opposé, opérateur unaire préfixé noté $-$, multiplication, opérateur binaire noté $*$, division, opérateur binaire infixé noté $/$, et exponentiation, opérateur binaire noté $\uparrow$. Les opérateurs binaires associent tous à gauche sauf l'exponentiation qui associe à droite, et l'ordre de priorité est $\uparrow > -(opposé) > \{*, /\} > \{+, -(soustraction)\}$.

- Donner le bon parenthésage des mots $3 - 2 - 1$ et $2 \uparrow 3 \uparrow 4$.

Correction : Pour le premier mot, il faut appliquer la règle d'associativité gauche de la soustraction. La bonne lecture est donc $(3 - 2) - 1$. Pour lire le second mot il faut utiliser l'associativité droite de $\uparrow$. Cela donne $2 \uparrow (3 \uparrow 4)$.

- Donner le bon parenthésage du mot $x + -y \uparrow 2 - 3 * x$.

Correction : L'opérateur le plus prioritaire est $\uparrow$. La bonne lecture est donc $(x + (-(y \uparrow 2))) - (3 * x)$.

- Donner une grammaire non-ambiguë décrivant ce langage.

Correction : L'associativité à gauche d'un opérateur op peut être codée dans un grammaire à l'aide de règles de production de la forme $E \rightarrow E \text{ op } E'$ où le E' est soit un terminal soit un symbole non terminal pour la production des expressions ayant un opérateur principal de plus grande priorité que op .

La priorité à droite se code avec des règles de la forme $E \rightarrow E' \text{ op } E$ où le E' est soit un terminal soit un symbole non terminal pour la production des expressions ayant un opérateur principal de plus grande priorité que op .

Pour pouvoir engendrer toutes les expressions il ne faut pas oublier qu'une expression n'a pas forcément comme symbole principal l'opérateur de plus faible priorité.

$$\begin{aligned} E &\rightarrow E + E_1 \mid E - E_1 \mid E_1 \\ E_1 &\rightarrow E_1 * E_2 \mid E_1 / E_2 \mid E_2 \\ E_2 &\rightarrow -E_2 \mid E_3 \\ E_3 &\rightarrow \text{int} \mid \uparrow E_3 \mid \text{int} \end{aligned}$$

Exercice 4.5

1. Spécifier un type abstrait « arbre de syntaxe », avec ses fonctions de construction et d'accès.

Correction :

```
type abstrait << ads >> (pour <<arbre de syntaxe >> ):
```

```
constructeurs:
```

```
noeud : lettre x (sequence d'ads) -> ads
{ noeud(a,s) construit l'arbre de syntaxe dont la racine est la
lettre a et dont les sous-arbres sont les arbres de la sequence s }
```

```
acces:
```

```
racine : ads -> lettre
{ racine(a) retourne la racine de l'arbre a }
```

```
sous_arbres : ads -> sequence d'ads
{ sous_arbres(a) retourne la sequence des sous-arbres de a }
```

le type sequence est suppose connu, defini par :

```
type abstrait << sequence de <element> >>:
```

```
constantes:
```

```
seq_vide
{ seq_vide représente la sequence vide }
```

```
constructeurs:
```

```
cons_seq : element x sequence -> sequence
{ cons_seq(x,s) retourne la sequence formee de x suivi des elements
de s }
```

```
tests:
```

```
seq_est_vide : sequence -> boolean
{ seq_est_vide(s) retourne vrai si la sequence est vide }
```

```
acces:
```

```
debut : sequence -> element
{ debut(s) retourne le premier element de la sequence s (erreur si
s est vide ) }
```

```
reste : sequence -> sequence
{ reste(s) retourne la sequence formee de s exceptee le premier
element (erreur si s est vide ) }
```

2. Spécifier puis réaliser une fonction de parcours gauche et une fonction de parcours droit.

Correction :

```

fonction parcours_gauche:

specifications:

    parcours_gauche : ads -> sequence de lettres
    { parcours_gauche(a) retourne ... }

    parcours_gauche_sequence : sequence d'ads -> sequence de lettres
    { parcours_gauche_sequence(s) retourne ... }

realisations :

parcours_gauche(a):
    cons_seq(racine(a),parcours_gauche_sequence(sous_arbres(a))

parcours_gauche_sequence(s):
    si seq_est_vide(s)
        alors seq_vide
    sinon concat(parcours_gauche(debut(s)),parcours_gauche_sequence(reste(s)))

```

3. Coder ce type abstrait et ces fonctions de parcours en CAML puis en PASCAL.

Exercice 4.6

1. Construire l'arbre de dérivation du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire ambiguë des expressions arithmétiques, correspondant au parenthésage $((x + (2.5 \times 4)) + (y + z))$.
2. Construire l'arbre de dérivation correspondant à la dérivation gauche du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire désambiguïsée des expressions arithmétiques.

Exercice 4.7

Considérons le langage L_C des expressions conditionnelles, où l'on peut construire des conditionnelles avec partie *else* (dites closes) ou sans partie *else* (dites non closes). Par exemple, les expressions suivantes sont dans L_C (c représente une condition quelconque, et i représente une instruction quelconque):

```

if c then i
if c then i else i
if c then i else if c then i else i

```

1. Trouvez une grammaire pour L_C .

Correction : Deux solutions immédiates et raisonnables sont: $I \rightarrow \text{if } C \text{ then } I \text{ else } I \mid \text{if } C \text{ then } I \text{ et } I \rightarrow \text{if } C \text{ then } I \text{ Opt, } \text{Opt} \rightarrow \varepsilon \mid \text{else } I$.

2. Donnez un arbre de dérivation pour la troisième expression. Votre grammaire est-elle ambiguë?
3. Trouvez une grammaire non-ambiguë pour L_C . On souhaite implémenter la convention courante qui veut que le **else** soit rattaché au **if then** le plus proche. (Considérez deux types de règle de production, une pour les instructions closes, l'autre pour les instructions non closes.)

Correction : I_{clos} (resp. $I_{nonclos}$) engendre tous les expressions où le dernier conditionnel est clos (non clos).

$$\begin{array}{ll}
 I & \rightarrow I_{clos} \mid I_{nonclos} \\
 I_{clos} & \rightarrow \text{if } C \text{ then } I_{clos} \text{ else } I_{clos} \\
 I_{nonclos} & \rightarrow \text{if } C \text{ then } I_{clos} \text{ else } I_{nonclos} \\
 & \mid \text{if } C \text{ then } I
 \end{array}$$

4. Est-ce que vous conseillez la première grammaire pour la définition d'un langage de programmation ou bien la deuxième? Quels changements dans la syntaxe proposez-vous pour ne pas avoir à choisir.

Correction : Solutions proposées: Ajouter des parenthèses (*lisp*) ou *begin end* (?), ajouter un *fi* (*sh*) un *endif* (*csh*, *cpp*) à la fin de la construction, faire la partie *else* (peut-être avec une instruction vide) obligatoire (ML sur sa grammaire abstraite ne considère que les conditionnelles complètes. Une absence de **else** est prise comme un **else** () implicite).



Exercice 4.8

Le langage $\{a^n b^n c^n \mid n \in \mathbb{N}\}$ est-il engendré par une grammaire hors-contexte? On pourra s'inspirer de l'exercice 2.3.

Exercice 4.9

Un palindrome non trivial est un mot de longueur au moins deux sur l'alphabet $a - z$ qui se lit aussi bien de droite à gauche que de gauche à droite. Par exemple, *eve*, *anna*, *otto* sont des palindromes. Une suite-palindrome sera un mot sur ce même alphabet obtenu en concaténant plusieurs palindromes, comme par exemple *annaeveotto*.

1. Donner une grammaire pour les suites-palindrome.

Correction :

$$G = (\{a, \dots, z\}, \{S\}, S, \{S\}, \{S \rightarrow aSa \mid \dots \mid zSz \mid S \rightarrow aTa \mid \dots \mid zTz \mid SS \mid \varepsilon, T \rightarrow a \mid \dots \mid z\}$$

2. Dessiner les arbres de dérivation des mots *annaeveotto* et *aaaa*

Correction : un arbre pour *annaeveotto* et deux pour *aaaa*.

3. Est-ce que la grammaire est ambiguë?

Correction : Oui. Voir les dérivations du second mot ci-dessus.

4. Peut-on trouver une grammaire non ambiguë pour le langage?

Correction : Non, car le même mot peut avoir plusieurs découpages inclus les uns dans les autres, ou au contraire se chevauchant, et le choix d'un découpage minimal ou maximal ne peut se faire par une grammaire context-free (je pense).

Exercice 4.10

On souhaite engendrer tous les mouvements d'une machine à découper le tissu. Pour faciliter le travail, nous allons supposer que la pièce de tissu a une longueur et une hauteur infinie. Initialement la machine se trouve à l'angle inférieur gauche de la pièce de tissu.

1. À quelle découpe correspond le mot $\uparrow\uparrow\rightarrow\rightarrow\downarrow\downarrow$?

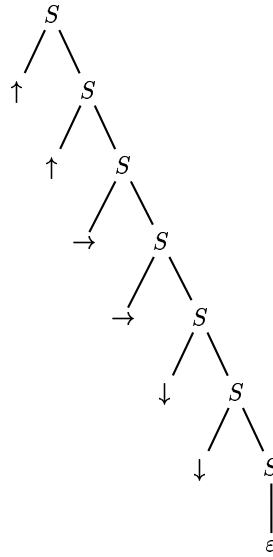
Correction : C'est un carré.

2. Trouver une grammaire qui permet d'engendrer toutes les découpes.

Correction : $S \rightarrow \downarrow S \mid \uparrow S \mid \leftarrow S \mid \rightarrow S \mid \varepsilon$.

3. Donnez un arbre de dérivation de l'exemple pour votre grammaire. Votre grammaire est-elle ambiguë?

Correction :



Non cette grammaire n'est pas ambiguë : une seule règle s'applique pour chaque début de mot.

4. Donnez une grammaire qui ne permet de découper que des rectangles.

$$S \rightarrow \uparrow S \downarrow \mid U$$

Correction : $U \rightarrow \varepsilon \mid \rightarrow U$

5. Donnez une grammaire qui permet de ne découper que des carrés.

Correction : On veut donc des mots de la forme $\uparrow^n \rightarrow^n \downarrow^n$. ce n'est pas possible, les grammaires non contextuelles ne permettent pas de tels mots. Pour la correction de cette question, il va falloir réfléchir et trouver un contre-exemple à toutes leurs propositions.

6. On suppose que le mouvement $\leftarrow$ n'est plus possible. Donner une grammaire non-ambiguë qui engendre toutes les découpes possibles sans sortir de la pièce de tissu.

Correction : $S \rightarrow S \uparrow S \downarrow S \mid \varepsilon \mid \rightarrow$

7. Si maintenant le mouvement $\leftarrow$ est de nouveau possible, pouvez-vous donner une grammaire qui engendre toutes les découpes qui ne sortent pas de la pièce de tissu?

Correction : Non.

Exercice 4.11

On considère le langage des expressions arithmétiques avec addition, soustraction, opposé, multiplication, division, et exponentiation (notée $\uparrow$) :

$$\begin{aligned} E &\rightarrow E + E \mid E - E \mid -E \mid E \times E \mid E / E \mid E \uparrow E \mid (E) \mid Int \\ Int &\rightarrow [0 - 9]^+ \end{aligned}$$

Cette grammaire étant ambiguë, on adopte les règles de désambiguïsation suivantes : les opérateurs binaires associent à gauche sauf l'exponentiation qui associe à droite, et l'ordre de priorité décroissant est l'exponentiation, puis l'opposé, puis la multiplication et la division, puis l'addition et la soustraction.

1. Donner le parenthésage implicite correct du mot $1 - 3 \times 4 + -5 \uparrow 2$.

Correction : $(1 - (3 \times 4)) + (-(5 \uparrow 2))$.

2. Définir une grammaire non-ambiguë décrivant ce langage, respectant les règles de désambiguïsation. On utilisera un nouveau non-terminal par niveau de priorité.

Correction :

$$\begin{aligned}
 E &\rightarrow E + F \mid E - F \mid F \\
 F &\rightarrow F \times G \mid F / G \mid G \\
 G &\rightarrow -G \mid H \\
 H &\rightarrow I \uparrow H \mid I \\
 I &\rightarrow (E) \mid Int
 \end{aligned}$$

3. Associer à chaque règle de cette grammaire un calcul permettant de déterminer le nombre d'opérateurs de l'expression reconnue.

Correction :

$$\begin{aligned}
 E &\rightarrow E + F & \{E_1 + F_1 + 1\} \\
 E &\rightarrow E - F & \{E_1 + F_1 + 1\} \\
 E &\rightarrow F & \{F_1\} \\
 F &\rightarrow F \times G & \{F_1 + G_1 + 1\} \\
 F &\rightarrow F / G & \{F_1 + G_1 + 1\} \\
 F &\rightarrow G & \{G_1\} \\
 G &\rightarrow -G & \{G_1 + 1\} \\
 G &\rightarrow H & \{H_1\} \\
 H &\rightarrow I \uparrow H & \{I_1 + H_1 + 1\} \\
 H &\rightarrow I & \{I_1\} \\
 I &\rightarrow (E) & \{E_1\} \\
 I &\rightarrow Int & \{0\}
 \end{aligned}$$

4. Associer à chaque règle de la grammaire non ambiguë un calcul permettant d'évaluer l'expression reconnue.

Correction :

$$\begin{aligned}
 E &\rightarrow E + F & \{E_1 + F_1\} \\
 E &\rightarrow E - F & \{E_1 - F_1\} \\
 E &\rightarrow F & \{F_1\} \\
 F &\rightarrow F \times G & \{F_1 G_1\} \\
 F &\rightarrow F / G & \{F_1 / G_1\} \\
 F &\rightarrow G & \{G_1\} \\
 G &\rightarrow -G & \{-G_1\} \\
 G &\rightarrow H & \{H_1\} \\
 H &\rightarrow I \uparrow H & \{I_1^{H_1}\} \\
 H &\rightarrow I & \{I_1\} \\
 I &\rightarrow (E) & \{E_1\} \\
 I &\rightarrow Int & \{val(Int_1)\}
 \end{aligned}$$

Exercice 4.12

On se pose le problème d'afficher un texte sur un terminal dont le nombre de colonnes est fixé C , de manière à ce qu'aucun mot ne soit à cheval sur deux lignes. On veut aussi « compresser » le texte en supprimant les occurrences consécutives des séparateurs (caractères blancs et passages à la ligne, notés respectivement `_` et `/` ci-dessous). Ainsi un texte source peut être :

```
qui__veut__voyager_loin__/_menage_sa__monture
```

La sortie d'un tel texte sur 20 colonnes doit être

```
qui_veut_voyager/loin_menage_sa/monture
```

Le texte source respecte la grammaire suivante :

$$\begin{aligned}
 \textit{Phrase} &\rightarrow \textit{Mot} \mid \textit{Phrase Sep Phrase} \\
 \textit{Mot} &\rightarrow \textit{Car} \mid \textit{Mot Mot} \\
 \textit{Car} &\rightarrow \textit{a} \mid \dots \mid \textit{z} \\
 \textit{Sep} &\rightarrow \textit{-} \mid \textit{/} \mid \textit{Sep Sep}
 \end{aligned}$$

1. Montrer que cette grammaire est ambiguë.
2. Définir une grammaire non-ambiguë engendrant le même langage.
3. Associer à chaque règle de la grammaire non-ambiguë obtenue un calcul qui permette de calculer le texte formaté correctement. Les valeurs calculées seront des enregistrements $\{long : int; texte : string\}$ où *texte* est le texte formaté et *long* donne la longueur de la chaîne *texte* après le dernier retour à la ligne.

Correction :

Pour avoir un calcul simple, il faut faire des règles récursives gauches.

$$\begin{aligned}
 M &\rightarrow c && \{(long = 1, texte = val(c))\} \\
 M &\rightarrow Mc && \{(long = M_1.long + 1, texte = M_1.texte^val(c))\} \\
 P &\rightarrow M && \{M_1\} \\
 P &\rightarrow PSM && \{si P_1.long + 1 + M_1.long > C \\
 &&& \quad \text{alors } (long = M_1.long, texte = P_1.texte^{"^"}M_1.texte) \\
 &&& \quad \text{sinon } (long = P_1.long + 1 + M_1.long, texte = P_1.texte^{"^"}M_1.texte)\} \\
 S &\rightarrow - && \{\} \\
 S &\rightarrow / && \{\} \\
 S &\rightarrow S_ && \{\} \\
 S &\rightarrow S/ && \{\}
 \end{aligned}$$

4. Déterminer la valeur calculée par cette grammaire pour l'exemple ci-dessus, au cours d'une analyse ascendante gauche.

Correction : On calcule au cours d'une analyse ascendante du texte

qui__veut__voyager_loin__/_menage_sa__monture

$M \rightarrow C$

M ui__veut__voyager_loin__/_menage_sa__monture

où $M \Rightarrow (long = 1, texte = "q")$

$M \rightarrow MC$

M i__veut__voyager_loin__/_menage_sa__monture

x où $M \Rightarrow (long = 2, texte = "qu")$

$M \rightarrow MC$

M __veut__voyager_loin__/_menage_sa__monture

où $M \Rightarrow (long = 3, texte = "qui")$

$P \rightarrow M$

`P __veut__voyager_loin__/_menage_sa__monture`

où $P \Rightarrow (long = 3, texte = "qui")$

$S \rightarrow _$

`P S __veut__voyager_loin__/_menage_sa__monture`

où $P \Rightarrow (long = 3, texte = "qui")$

$S \rightarrow S_$

`P S _veut__voyager_loin__/_menage_sa__monture`

où $P \Rightarrow (long = 3, texte = "qui")$

etc.

`P S M __voyager_loin__/_menage_sa__monture`

où $P \Rightarrow (long = 3, texte = "qui")$ et $M \Rightarrow (long = 4, texte = "veut")$

$P \rightarrow PSM$

`P __voyager_loin__/_menage_sa__monture`

où $P \Rightarrow (long = 8, texte = "qui_veut")$

etc.

`P S M __/_menage_sa__monture`

où $P \Rightarrow (long = 16, texte = "qui_veut_voyager")$ et $M \Rightarrow (long = 4, texte = "loin")$

comme $16 + 1 + 4 > 20$:

`P __/_menage_sa__monture`

où $P \Rightarrow (long = 4, texte = "qui_veut_voyager/loin")$

etc.

5. Modifier les calculs de la grammaire de manière à ce que lorsqu'une séquence de séparateurs contient au moins un retour à la ligne, alors cela force le passage à la ligne dans le texte formaté. Sur l'exemple, cela doit donner :

`qui_veut_voyager/loin/menage_sa_monture`

car le texte source contient un retour à la ligne entre « loin » et « menage ». (On pourra donner une valeur booléenne aux non-terminaux *Sep*.)

Correction :

$$\begin{aligned}
 S &\rightarrow _ && \{false\} \\
 S &\rightarrow / && \{true\} \\
 S &\rightarrow S_ && \{S_1\} \\
 S &\rightarrow S/ && \{true\} \\
 M &\rightarrow c && \{(long = 1, texte = val(c))\} \\
 M &\rightarrow Mc && \{(long = M_1.long + 1, texte = M_1.texte^val(c))\} \\
 P &\rightarrow M && \{M_1\} \\
 P &\rightarrow PSM && \{si\ S_1\ ou\ P_1.long + 1 + M_1.long > C \\
 &&& \quad \text{alors } (long = M_1.long, texte = P_1.texte^{"^"}M_1.texte) \\
 &&& \quad \text{sinon } (long = P_1.long + 1 + M_1.long, texte = P_1.texte^{"_"}M_1.texte)\}
 \end{aligned}$$

6. (question optionnelle) Comment peut-on faire pour « justifier » le texte c'est-à-dire répartir des blancs entre les mots pour que toutes les lignes soient pleines? (On demande de ne pas utiliser de fonction auxiliaire qui reparcourerait le texte à afficher.)

Exercice 4.13

On veut piloter un traceur graphique au moyen de mots sur l'alphabet $\{\uparrow, \downarrow, \rightarrow, \leftarrow\}$. Le langage accepté est donné par la grammaire suivante :

$$L \rightarrow LL \mid \uparrow \mid \downarrow \mid \rightarrow \mid \leftarrow \mid \varepsilon$$

Le traceur est situé initialement en position (0,0). Chaque mouvement le déplace de 1mm.

1. Donnez une grammaire non ambiguë pour ce langage.

Correction : il existe une variante recursive droite :

$$L \rightarrow \varepsilon \mid \uparrow L \mid \downarrow L \mid \leftarrow L \mid \rightarrow L$$

et une récursive gauche :

$$L \rightarrow \varepsilon \mid L \uparrow \mid L \downarrow \mid L \leftarrow \mid L \rightarrow$$

2. Donnez les expressions associées à cette grammaire pour calculer la position du pointeur en fin de lecture du mot. Ces expressions seront de type enregistrement $\{x : int; y : int\}$.

Correction : Il vaut mieux prendre la variante récursive gauche :

$$\begin{aligned}
 L &\rightarrow \varepsilon && \{x = 0 \ ; \ y = 0\} \\
 L &\rightarrow L \uparrow && \{x = L_1.x \ ; \ y = L_1.y + 1\} \\
 L &\rightarrow L \downarrow && \{x = L_1.x \ ; \ y = L_1.y - 1\} \\
 L &\rightarrow L \leftarrow && \{x = L_1.x - 1 \ ; \ y = L_1.y\} \\
 L &\rightarrow L \rightarrow && \{x = L_1.x + 1 \ ; \ y = L_1.y\}
 \end{aligned}$$

Attention à bien respecter la notation N_k pour désigner la valeur associée à la k -ième occurrence du non-terminal N dans le membre droit d'une règle.

3. Déterminer les calculs effectués lors de l'analyse ascendante gauche du mot $\uparrow \rightarrow \uparrow \leftarrow$. On dessinera au préalable l'arbre de dérivation et l'on effectuera les calculs en remontant dans cet arbre, puis on montrera comment ces calculs sont effectués par l'analyse ascendante en écrivant à chaque étape les racines des arbres de dérivation intermédiaires.

Correction : Il faut profiter de cette question pour bien montrer l'analyse ascendante gauche. Le dessin de l'arbre est laissé au lecteur, on donne ici simplement les racines des arbres lors de l'analyse.

$\uparrow \rightarrow \uparrow \leftarrow$

Réduction par la règle $L \rightarrow \varepsilon$

$L \uparrow \rightarrow \uparrow \leftarrow$

avec $L \Rightarrow \{x = 0; y = 0\}$

Réduction par la règle $L \rightarrow L \uparrow$

$L \rightarrow \uparrow \leftarrow$

avec $L \Rightarrow \{x = 0; y = 1\}$

Réduction par la règle $L \rightarrow L \rightarrow$

$L \uparrow \leftarrow$

avec $L \Rightarrow \{x = 1; y = 1\}$

Réduction par la règle $L \rightarrow L \uparrow$

$L \leftarrow$

avec $L \Rightarrow \{x = 1; y = 2\}$

Réduction par la règle $L \rightarrow L \leftarrow$

L

avec $L \Rightarrow \{x = 0; y = 2\}$

Exercice 4.14

Les deux grammaires suivantes engendrent le langage L_s des séquences d'identificateurs :

$S \rightarrow \varepsilon \mid Id S$

$S \rightarrow \varepsilon \mid S Id$

1. Laquelle de ces deux grammaires permet-elle une analyse descendante du langage L_s ?
2. Supposons qu'on veuille réaliser une analyse ascendante pour le langage L_s avec un outil comme YACC. Laquelle des deux grammaires est préférable si on veut économiser la consommation de mémoire?

Exercice 4.15

Dans le fichier YACC des expressions arithmétiques, ajouter des expressions pour calculer :

1. le nombre d'opérateurs ;
2. l'ensemble des identificateurs utilisés ;
3. l'arbre de dérivation.

Dans chaque cas, spécifier d'abord les opérations auxiliaires nécessaires.

Chapitre 5

Syntaxe abstraite des langages

L'objectif de ce chapitre est de définir une structure de données permettant de représenter le résultat de l'analyse syntaxique sous une forme adaptée à l'analyse sémantique.

Le type et la signification d'une phrase d'un langage sont calculés à partir de son arbre de syntaxe. Mais la syntaxe (dite *concrète*) d'un langage est souvent encombrée de complications nécessaires à l'analyse syntaxique (visant en particulier à traiter les problèmes d'ambiguïté), mais nuisibles à la compréhension, donc à la définition de sa sémantique. Cela rend nécessaire la définition d'un niveau intermédiaire, celui de la *syntaxe abstraite*, dont le but est d'éliminer de l'arbre de dérivation les détails devenus inutiles à ce stade. Plus précisément, la construction d'un *arbre de syntaxe abstraite* a trois objectifs :

- éviter les ambiguïtés inhérentes aux grammaires en construisant une structure nécessairement non-ambiguë ;
- éliminer les complications inutiles introduites dans les arbres de dérivations par l'abondance de mots clés et de symboles de ponctuation ;
- réunir en une seule notion les arbres de dérivation décrivant la structure d'une phrase du langage et les valeurs associées aux tokens de cette phrase.

Dans le cas des expressions arithmétiques, les expressions de syntaxe abstraite que l'on souhaite obtenir sont de la forme $+(Cte, \times(Cte, Cte))$. Ces expressions se prêtent en effet particulièrement bien à une évaluation, car elles correspondent à des arbres dont les nœuds sont étiquetés par des opérations arithmétiques. Une évaluation revient donc à propager les calculs d'opérations arithmétiques dans l'arbre.

5.1 Grammaires abstraites

Définition 5.1 Soit T un ensemble (fini ou infini) appelé ensemble des étiquettes abstraites et soit V_t le vocabulaire formé par T augmenté des parenthèses et de la virgule. Une grammaire G sur le vocabulaire terminal V_t est dite abstraite si

1. toutes les règles sont de la forme $N \rightarrow t$ ou bien $N \rightarrow t(N_1, \dots, N_n)$ ($n \geq 1$) où les N_i sont des non-terminaux et $t \in T$;
2. une étiquette t donnée n'apparaît qu'une seule fois dans l'ensemble des règles.

Un mot engendré par une grammaire abstraite s'appelle une expression de syntaxe abstraite ou, plus simplement, expression abstraite.

Par exemple, si $T = \{+, \times, \text{Cte}\}$ alors la grammaire

$$E \rightarrow +(E, E) \mid \times(E, E) \mid \text{Cte}$$

sera notre grammaire abstraite des expressions arithmétiques formées d'additions, de multiplications et de constantes.

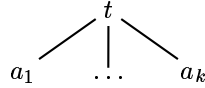
Grâce à la propriété qu'une étiquette t donnée n'apparaît qu'une seule fois dans l'ensemble des règles, on a la propriété fondamentale suivante :

Proposition 5.2 *Une grammaire abstraite est non-ambiguë.*

Notre premier objectif est donc atteint.

5.2 Arbres de syntaxe abstraite

Nous allons maintenant donner une représentation arborescente des expressions abstraites. Cela est possible grâce à la forme $N \rightarrow t(N_1, \dots, N_k)$ des règles d'une grammaire abstraite. Une expression engendrée par le non-terminal N de cette règle sera représentée par l'arbre



où a_i est, récursivement, l'arbre associé à l'expression abstraite e_i engendrée par le non-terminal N_i . On dira alors que cet arbre est de *sorte* N . D'où la définition :

Définition 5.3 *Étant donnée une grammaire abstraite $G = (V_t, V_n, S, R)$, un arbre de syntaxe abstraite de sorte $N \in V_n$ (ou plus simplement arbre abstrait) est un arbre A dont les nœuds sont étiquetés par des symboles $t \in T$, tel que :*

- soit il existe une règle $N \rightarrow t \in R$ et A est réduit à une feuille étiquetée par t ;
- soit il existe une règle $N \rightarrow t(N_1, \dots, N_n)$ telle que
 - (i) la racine de A soit étiquetée par t ;
 - (ii) A ait exactement n fils qui soient, récursivement, des arbres de syntaxe abstraite de sortes respectives $N_1, \dots, N_n$.

Pour la grammaire abstraite précédente, on a par exemple les arbres de syntaxe abstraite suivants :



Il existe une correspondance biunivoque entre ces arbres et les mots reconnus par la grammaire abstraite : le parcours gauche d'un arbre de syntaxe abstraite permet de construire le mot correspondant. Ainsi, les exemples d'arbres précédents correspondent aux mots

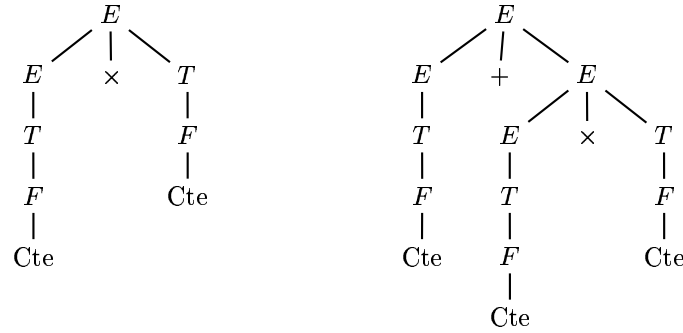
$$\times(\text{Cte}, \text{Cte})$$

$$+(\text{Cte}, \times(\text{Cte}, \text{Cte}))$$

Réciproquement, on associe à chaque mot engendré par la grammaire abstraite un arbre de syntaxe abstraite par une construction analogue à la construction récursive élaborée au paragraphe 4.2 :

- à un mot engendré par une règle $N \rightarrow t$ on associe l'arbre réduit à une feuille étiquetée par t ;
- à un mot engendré par une règle $N \rightarrow t(N_1, \dots, N_n)$ on associe l'arbre dont la racine est étiquetée par t et dont les fils sont, dans cet ordre, les arbres associés aux mots engendrés par $N_1, \dots, N_n$.

Si l'on compare les arbres de syntaxe abstraite des mots exemples précédents avec leurs arbres de dérivation dans la grammaire concrète non-ambiguë du langage des expressions arithmétiques :



on constate que l'on a atteint notre deuxième objectif: les non-terminaux introduits pour désambigüiser ont été éliminés. Remarquons aussi que si l'expression originale contenait des parenthèses, celles-ci ne sont plus présentes dans l'arbre: c'est la structure de l'arbre qui donne le parenthésage.

5.3 Représentation des valeurs des tokens dans l'arbre

Il nous reste à représenter les valeurs des tokens dans l'arbre de syntaxe abstraite lui-même. Pour ce faire, il n'est pas nécessaire de changer la notion d'arbre: il suffit d'utiliser la possibilité d'avoir un ensemble infini d'étiquettes abstraites. Cet ensemble sera donc en fait un langage régulier (les entiers, les chaînes de caractères, etc.) aussi cet ensemble d'étiquettes abstraites sera souvent informellement donné en extension avec des points de suspension, ou encore sous forme d'une expression rationnelle.

Ainsi, pour représenter les expressions arithmétiques avec les valeurs des constantes, il suffit de prendre $T = \{+, \times, \text{Cte}\} \cup \mathbb{N}$ et la grammaire abstraite suivante :

$$\begin{aligned} E &\rightarrow +(E, E) \mid \times(E, E) \mid \text{Cte}(\text{Nat}) \\ \text{Nat} &\rightarrow (0 \mid 1 \mid \dots \mid 9)^+ \end{aligned}$$

Le langage défini par cette grammaire contiendra donc par exemple $\times(\text{Cte}(1), \text{Cte}(2))$ et $+(\text{Cte}(1), \times(\text{Cte}(2), \text{Cte}(3)))$.

Pour résumer, la figure 5.1 montre un tableau donnant deux exemples d'expressions arithmétiques, accompagnées de leur arbre de dérivation et leur arbre abstrait associé, où l'on constate à nouveau la simplicité et la concision apportée par la notion d'arbre abstrait. Nous avons donc atteint notre dernier objectif.

syntaxe utilisateur	syntaxe concrète	arbre de dérivation	arbre de syntaxe abstraite
1+3*2	$Nat + Nat \times Nat$	The derivation tree for the expression 1+3*2 has root E. E has three children: E, +, and T. The left E has child T, which has child F, which has child Nat₁. The middle T has child F, which has child Nat₂. The right T has three children: ×, F, and Nat₃. $\begin{aligned} \text{val}(Nat_1) &= 1 \\ \text{val}(Nat_2) &= 3 \\ \text{val}(Nat_3) &= 2 \end{aligned}$	The abstract syntax tree for 1+3*2 has root +. The root has three children: Cte, ×, and Cte. The left Cte has child 1. The middle Cte has child 3. The right Cte has child 2.
(1+3)*2	$(Nat + Nat) \times Nat$	The derivation tree for the expression (1+3)*2 has root E. E has child T, which has three children: T, ×, and F. The left T has child F, which has three children: (, E, and). The inner E has three children: E, +, and T. The left E has child T, which has child F, which has child Nat₁. The middle T has child F, which has child Nat₂. The right T has child F, which has child Nat₃. $\begin{aligned} \text{val}(Nat_1) &= 1 \\ \text{val}(Nat_2) &= 3 \\ \text{val}(Nat_3) &= 2 \end{aligned}$	The abstract syntax tree for (1+3)*2 has root ×. The root has three children: +, Cte, and Cte. The left + has three children: Cte, Cte, and Cte. The leftmost Cte has child 1. The middle Cte has child 3. The rightmost Cte has child 2.

FIG. 5.1 – Comparaison des syntaxes abstraites et concrètes

5.4 Calcul des arbres de syntaxe abstraite

Le dernier point qui reste à traiter dans ce chapitre est la façon d'associer à chaque mot d'un langage son arbre abstrait. Cela se fait facilement dans le cadre du logiciel YACC comme expliqué au paragraphe 4.4. Ainsi, si l'on poursuit notre exemple des expressions arithmétiques, on aura :

$$\begin{array}{lll}
 E & \rightarrow & E + T \quad \{+(E_1, T_1)\} \\
 E & \rightarrow & T \quad \{T_1\} \\
 T & \rightarrow & T \times F \quad \{\times(T_1, F_1)\} \\
 T & \rightarrow & F \quad \{F_1\} \\
 F & \rightarrow & Nat \quad \{\text{Cte}(\text{val}(Nat_1))\} \\
 F & \rightarrow & (E) \quad \{E_1\}
 \end{array}$$

Noter l'utilisation de la fonction `val` pour accéder à l'unité lexicale représentée par le token `Nat`.

5.5 Codages des arbres de syntaxe abstraite

Nous terminons ce chapitre par quelques indications sur la méthode que l'on peut utiliser pour programmer ces arbres de syntaxe abstraite.

Dans le langage CAML, ce codage est immédiat en utilisant des types sommes. Ainsi, on définira le type des expressions arithmétiques par :

```

type expr =
  Cte of int
| Plus of expr * expr
| Mult of expr * expr

```

D'une manière générale, il y aura autant de types à définir que de sortes d'arbres de syntaxe abstraite, et il y aura autant de constructeurs que de règles de grammaire abstraite. Les ensembles infinis d'étiquettes abstraites comme les entiers ou les identificateurs seront codés simplement par le type de base correspondant : `int` ou `string`. Cela est notre seule façon de coder des ensembles infinis d'étiquettes abstraites. Nous imposerons donc que tout ensemble infini d'étiquettes abstraites corresponde à l'un des types de base parmi `int`, `string`, et `real`.

Dans un langage comme PASCAL ou C n'offrant pas cette possibilité des types sommes et d'allocation automatique de mémoire, il est nécessaire de coder ces arbres à l'aide de pointeurs sur des enregistrements, contenant un champ pour l'étiquette abstraite et des champs pour les sous-arbres. Dans un langage avec objets comme C++ ou JAVA, on aura avantage à coder les arbres par des classes, la construction et la destruction de ces arbres en seront facilitées.

5.6 Exercices

Exercice 5.1

1. Écrire une grammaire abstraite pour les expressions booléennes formées avec les opérations « ou », « et » et « non », et les constantes « vrai » et « faux ».
2. Donner une grammaire concrète non-ambiguë pour le langage des expressions booléennes avec toutes les opérations ci-dessus, et donner les règles de calcul de la syntaxe abstraite associées.
3. Donner un type CAML représentant cette syntaxe abstraite.

4. Écrire une fonction CAML permettant de transformer une chaîne de caractères représentant une expression booléenne en une expression abstraite, en utilisant l'outil CAMLYACC.

Exercice 5.2

1. Écrire une grammaire abstraite pour les expressions arithmétiques sur l'addition, la multiplication, la soustraction, la division, l'opposé, l'exponentiation, les constantes entières ou réelles et les variables. (On ne distinguera pas les constantes entières, on les codera par des réels.)
2. Donner une grammaire concrète non-ambiguë pour le langage des expressions ci-dessus, et donner les règles de calcul de la syntaxe abstraite associées.
3. Donner un type CAML représentant cette syntaxe abstraite.
4. Écrire une fonction CAML permettant de transformer une chaîne de caractères représentant une expression arithmétique en une expression abstraite, en utilisant l'outil CAMLYACC.

Deuxième partie

Sémantique

Chapitre 1

Généralités sur la sémantique des langages

Décrire, implémenter, et même utiliser un langage de programmation nécessite une description précise et complète de la signification, on dit la *sémantique*, des différentes phrases que l'on peut écrire dans le langage.

Nous procéderons toujours en trois temps pour définir la sémantique d'un langage: définition de la syntaxe abstraite dans un premier temps; définition de la fonction de typage dans un second temps, qui est une application des expressions de syntaxe abstraite dans l'ensemble des types; définition de la fonction d'évaluation dans un troisième temps, qui est une application des expressions de syntaxe abstraite dans un ensemble de valeurs indexé par l'ensemble des types. Cette structuration de la sémantique fait clairement apparaître que le typage ne dépend pas de l'évaluation. On dit que le typage est statique. Par contre, l'évaluation peut dépendre du typage. Dans le cas où l'évaluation ne dépend pas du typage, typage et évaluation calculeront à partir de la même expression de syntaxe abstraite. Le typage ne sera alors qu'un contrôle, oublié dès qu'il a réussi. Mais le typage peut également transformer l'expression de syntaxe abstraite afin d'y inclure des informations de type utilisées lors de l'évaluation.

Décrire la sémantique d'un langage de programmation est une œuvre complexe, et le formalisme utilisé peut dépendre du type de langage à décrire: fonctionnel, impératif, logique, à objets... Dans ce cours, nous allons utiliser un formalisme très général, celui des *règles sémantiques*, qui généralise la notion de grammaire, et même celle de grammaire avec actions présente dans YACC. L'idée consiste à calculer par récurrence sur des arbres à l'aide de règles de calcul appropriées.

1.1 Notion de jugement

Un jugement est une affirmation portant sur une expression abstraite. Avant d'en donner une définition précise, nous présentons plusieurs exemples simples de jugements :

- « l'expression $+(Cte(1), \times(Cte(2), Cte(3)))$ a pour valeur 7 » ;
- « l'expression $+(Cte(1), \times(Cte(2), Cte(3)))$ a pour type Int » ;
- « l'instruction $Write(+(Cte(3), Cte(4)))$ affiche 7 » ;
- « l'expression $+(Cte(3), Cte(5))$ a pour code machine $[PUSHI\ 3; PUSHI\ 5; ADD]$ ».

Un jugement affirme qu'à une certaine construction d'un langage (expression arithmétique, instruction, etc.) exprimée dans la syntaxe abstraite choisie, est associée une valeur d'un certain ensemble : un entier, un type, une chaîne de caractères, du code machine, etc.

Dans des cas plus complexes, l'affirmation pourra dépendre d'un contexte. Par exemple :

- « dans le contexte où x est une variable valant 4, l'expression $+(Cte(1), \times(Var(x), Cte(3)))$ a pour valeur 13 » ;
- « dans le contexte où x est une variable de type Real, l'expression $+(Cte(1), \times(Var(x), Cte(3)))$ a pour type Real » ;
- pour un compilateur devant engendrer du code machine à une adresse donnée, on pourra avoir des jugements du style « dans le contexte où l'adresse courante du code est 250 et x est une variable dont l'adresse est 100, l'expression $\times(Cte(3), Var(x))$ a pour code machine [250 PUSHI 3; 251 PUSH 100; 252 MUL] ».

Ces exemples conduisent à la définition suivante :

Définition 1.1 *Un jugement est un triplet $(C, a, v) \in D^c \times A \times D^v$ tel que :*

- (i) *A est un ensemble d'arbres de syntaxe abstraite, contenant l'arbre a de sorte N ;*
- (ii) *D^v est un ensemble appelé domaine de valeurs ; la valeur v appartient au sous-ensemble D_N^v des valeurs de sorte N .*
- (iii) *D^c est un ensemble appelé domaine de contextes ; la valeur C appartient au sous-ensemble D_N^c des valeurs de sorte N .*

On note habituellement le jugement $(C, a, v) \in D^c \times A \times D^v$ par « $C \vdash a \Rightarrow v$ ». Dans les cas simples où le domaine de contexte est absent (c.-à-d. plus précisément $D^c = \{\emptyset\}$, le seul contexte possible est le contexte vide), on notera en général simplement $a \Rightarrow v$. Ainsi :

- $+(Cte(1), \times(Cte(2), Cte(3))) \Rightarrow 7$ est le jugement « l'expression $1 + 2 \times 3$ a pour valeur 7 », ici $D^v = \mathbb{N}$ et $D^c = \{\emptyset\}$;
- $\{x \Rightarrow 4\} \vdash +(Cte(1), \times(Var(x), Cte(3))) \Rightarrow 13$ est le jugement « dans le contexte où x est une variable valant 4, l'expression $1 + x \times 3$ a pour valeur 13 », ici $D^v = \mathbb{N}$ et D^c est l'ensemble des *environnements d'évaluation* (cf. chapitre 3).

Lorsque le domaine de valeurs D^v est un ensemble de types, par exemple $D^v = \{\text{Int}, \text{Bool}, \text{Real}\}$, nous utiliserons une variante classique de cette notation, $C \vdash a : t$, où $t \in D^v$. On écrira donc :

- $+(Cte(1), \times(Cte(2), Cte(3))) : \text{Int}$ pour le jugement « l'expression $1 + 2 \times 3$ a pour type Int » ;
- $\{x : \text{Real}\} \vdash +(Cte(1), \times(Var(x), Cte(3))) : \text{Real}$ pour le jugement « dans le contexte où x est une variable de type Real, l'expression $1 + x \times 3$ a pour type Real ».

1.2 Règles sémantiques

La définition précédente autorise l'écriture de jugements vrais aussi bien que de jugements faux :

- « l'expression $+(Cte(1), \times(Cte(2), Cte(3)))$ a pour valeur 12 » ;
- « l'expression $+(Cte(1), \times(Cte(2), Cte(3)))$ a pour type Bool » ;

sont également des jugements¹

Étant donné un langage et une certaine catégorie de jugements associés, l'étape suivante dans la définition de la sémantique du langage consiste donc à déterminer quels sont les jugements vrais et les jugements faux, ce que nous allons faire par récurrence sur les arbres de syntaxe abstraite. Nous allons utiliser pour cela un outil classique permettant de construire des récurrences sur des arbres, appelé *règles sémantiques*. Parmi ces règles, certaines correspondent au cas de base de la récurrence, donc s'appliquent aux feuilles de l'arbre de syntaxe abstraite. Ce sont les *axiomes*, qui affirment qu'un certain jugement (ou ensemble de jugements) est vrai. Par exemple :

- si n est une constante entière, alors le jugement $\text{Cte}(n) \Rightarrow n$ (c.-à-d. « l'expression réduite à la constante n a pour valeur n ») est vrai ;
- si n est une constante entière, alors le jugement $\text{Cte}(n) : \text{Int}$ (c.-à-d. « l'expression réduite à la constante n a pour type Int ») est vrai ;
- dans un contexte C où x est une variable de type Bool , le jugement $C \vdash \text{Var}(x) : \text{Bool}$ (c.-à-d « l'expression réduite à l'identificateur x a pour type Bool ») est vrai.

Les règles proprement dites correspondent au cas général de la récurrence. Elles affirment la vérité d'un jugement qui constitue la *conclusion* de la règle sous l'hypothèse (de récurrence) de la vérité d'un ou plusieurs jugements appartenant aux *prémisses* de la règle. Par exemple :

- si le jugement $e_1 : \text{Int}$ (c.-à-d « l'expression e_1 a pour type Int ») est vrai et si le jugement $e_2 : \text{Int}$ (c.-à-d « l'expression e_2 a pour type Int ») est vrai alors le jugement $+(e_1, e_2) : \text{Int}$ (c.-à-d « l'expression $+(e_1, e_2)$ a pour type Int ») est également vrai.

Ces considérations nous conduisent à la définition formelle suivante :

Définition 1.2 Une règle sémantique est un k -uplet de jugements, avec $k \geq 1$. Si $k = 1$, la règle est appelée un *axiome*. Dans le cas d'une règle $(J_1, \dots, J_k)$ avec $k \geq 2$, $J_1, \dots, J_{k-1}$ constituent les prémisses de la règle et J_k en est la conclusion.

La notation usuelle pour la règle $(J_1, \dots, J_k)$ consiste à écrire une fraction avec les prémisses au-dessus du trait et la conclusion en-dessous :

$$\frac{J_1, \dots, J_{k-1}}{J_k}$$

Lorsque l'ensemble des prémisses est vide, c'est-à-dire dans le cas d'un axiome, rien ne figure au dessus du trait de fraction.

Imaginons par exemple que l'on veuille décrire l'évaluation d'une expression arithmétique à l'aide de règles sémantiques. Nous sommes amenés à écrire tous les axiomes :

$$\frac{}{\text{Cte}(0) \Rightarrow 0} \quad \frac{}{\text{Cte}(1) \Rightarrow 1} \quad \text{etc.}$$

ce qui, par une notation mathématique habituelle, se note de manière finie par le *schéma d'axiomes*

$$\text{pour tout } n \in \mathbb{N}: \quad \frac{}{\text{Cte}(n) \Rightarrow n}$$

1. Nous sommes ici au cœur d'une discussion philosophique chère aux logiciens : la vérité préexiste-t-elle ou non ? Pour les uns, les règles sémantiques telles que nous allons les introduire servent à vérifier si un jugement appartient à la catégorie, connue a-priori, des jugements vrais ; Pour les autres, la vérité est donnée par les règles sémantiques elles-mêmes.

De la même façon, on a les schémas de règles suivantes pour l'addition et la multiplication d'entiers :

$$\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{+(e_1, e_2) \Rightarrow v_1 +_{int} v_2}$$

et

$$\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{\times(e_1, e_2) \Rightarrow v_1 \times_{int} v_2}$$

En effet, e_1 n'est pas une expression abstraite, de même que v_1 n'est pas une valeur. Les schémas de règles ci-dessus désignent en fait l'ensemble infini de toutes les règles obtenues en remplaçant e_1 et e_2 par des expressions abstraites, ainsi que v_1 et v_2 par des valeurs entières. Nous ne reviendrons plus sur cette subtile nuance par la suite.

Notons la différence entre le symbole $+$ de l'expression de syntaxe abstraite, qui est un symbole du vocabulaire de la grammaire abstraite, et le symbole $+_{int}$ que nous avons volontairement indexé par le mot *int*, qui est l'opération mathématique habituelle sur l'ensemble des valeurs entières. Par la suite, la paresse pourra nous inciter à omettre cet indice, sans pour cela créer la moindre ambiguïté.

Notons également la ressemblance entre les jugements ci-dessus et l'évaluation des arbres syntaxiques telle qu'on peut la faire dans YACC comme expliqué au paragraphe 4.4. En fait, YACC permet, grâce aux expressions accompagnant la description de la syntaxe concrète d'un langage, de spécifier des jugements mais portant dans ce cas sur les arbres de syntaxe concrète.

1.3 Notion de dérivation d'un jugement

Les règles sémantiques vont nous servir à calculer des jugements vrais en partant des axiomes et en appliquant les règles un nombre quelconque de fois :

Définition 1.3 Une dérivation sémantique est une suite finie de jugements $J_0, \dots, J_n$ telles que pour tout $i \in [0, \dots, n]$:

- soit J_i est un axiome ;
- soit il existe une règle

$$\frac{J_{i_1}, \dots, J_{i_k}}{J_i}$$

avec $i_1, \dots, i_k \in [0, \dots, i-1]$. On dira encore qu'il s'agit d'une dérivation du jugement J_n .

Souvent, on omettra l'adjectif sémantique quand on parlera clairement de dérivation de jugement et pas d'autres types de dérivation. Voici par exemple une dérivation du jugement $+(Cte(1), \times(Cte(2), Cte(3))) \Rightarrow 7$:

$J_0 =$	$Cte(1) \Rightarrow 1$	axiome
$J_1 =$	$Cte(2) \Rightarrow 2$	axiome
$J_2 =$	$Cte(3) \Rightarrow 3$	axiome
$J_3 =$	$\times(Cte(2), Cte(3)) \Rightarrow 6$	règle de multiplication de prémisses J_1 et J_2
$J_4 =$	$+(Cte(1), \times(Cte(2), Cte(3))) \Rightarrow 7$	règle d'addition de prémisses J_0 et J_3

Cette notation linéaire des dérivations est ni très lisible ni très pratique. Une notation beaucoup plus agréable consiste à reproduire la structure de la dérivation. On obtient alors

un *arbre de dérivation d'un jugement* qui mime l'arbre de syntaxe abstraite. Sur l'exemple ci-dessus cela donne

$$\begin{array}{c}
 +(Cte(1), \times(Cte(2), Cte(3))) \Rightarrow 7 \\
 \swarrow \quad \searrow \\
 Cte(1) \Rightarrow 1 \quad \times(Cte(2), Cte(3)) \Rightarrow 6 \\
 \quad \swarrow \quad \searrow \\
 \quad Cte(2) \Rightarrow 2 \quad Cte(3) \Rightarrow 3
 \end{array}$$

Un arbre de dérivation est donc un arbre étiqueté par des jugements, dont chaque feuille est étiquetée par un axiome, et dont le jugement étiquetant un nœud interne est obtenu par une règle sémantique à partir des jugements étiquetant les fils de ce nœud. On appelle alors cet arbre *arbre de dérivation* du jugement étiquetant la racine.

Afin de reproduire l'écriture des règles sémantiques sous forme de fractions, l'habitude a été prise² de dessiner l'arbre de dérivation d'un jugement avec la racine en bas et les feuilles en haut. On dessine alors des traits de fraction en lieu et place des traits reliant habituellement un nœud et ses fils. Sur l'exemple, nous obtenons :

$$\frac{\frac{Cte(1) \Rightarrow 1}{\frac{\frac{Cte(2) \Rightarrow 2 \quad Cte(3) \Rightarrow 3}{\times(Cte(2), Cte(3)) \Rightarrow 6}}{+(Cte(1), \times(Cte(2), Cte(3))) \Rightarrow 7}}$$

1.4 Codage des règles sémantiques

Les règles sémantiques permettent de décrire la sémantique d'un langage. Mais elles constituent un langage de description très abstrait, qu'il nous faut maintenant implémenter sous la forme, par exemple, d'un interprète.

L'idée la plus naturelle serait de considérer que la sémantique d'un programme est calculée par une fonction qui prend un jugement en argument, et retourne une valeur booléenne indiquant si le jugement en question est vrai ou non :

sémantique : contexte $\times$ expression $\times$ valeur $\rightarrow$ booléen
 { sémantique(C, e, v) retourne vrai si le jugement $C \vdash e \Rightarrow v$ est vrai }

Mais ce n'est pas ce que l'on attend d'un interprète : on a besoin de déterminer la valeur d'une expression à l'aide d'une fonction dont la spécification est la suivante :

sémantique : contexte $\times$ expression $\rightarrow$ valeur
 { sémantique(C, e) retourne une valeur v telle que le jugement $C \vdash e \Rightarrow v$ soit vrai }

Bien entendu, une telle valeur n'existe pas toujours, auquel cas la fonction déclenchera un message d'erreur : ce sont les *erreurs de sémantique* : variable non déclarée, expression mal typée, etc. Il se peut également que plusieurs valeurs soient possibles. Nous veillerons toujours, dans ce cours, à éviter de tels systèmes de règles, dits non-déterministes.

Le codage de cette fonction sera directement inspiré de la formulation des règles sémantiques elles-mêmes : on utilisera une décomposition en cas suivant la forme de l'expression abstraite fournie en argument, et on codera les différents jugements figurant dans les prémisses d'une règle sous la forme d'appels récursifs. Par exemple, le jugement d'évaluation des expressions arithmétiques peut être codé en la fonction CAML suivante :

let rec valeur = fonction

2. Depuis trop longtemps pour en changer : la notion d'arbre de dérivation est née dans la communauté de logique mathématique. Même si cette communauté et celles de linguistique et d'informatique sont depuis longtemps conscientes des liens étroits qui existent entre leur disciplines, la formulation des concepts communs et leur notation sont souvent dissemblables d'une communauté à l'autre.

```

Cte(n) -> n                (* règle (Constante) *)
| Plus(e1,e2) ->
  let v1 = (valeur e1)      (* règle (Addition) *)
  and v2 = (valeur e2)
  in v1+v2
| Mult(e1,e2) ->
  let v1 = (valeur e1)      (* règle (Multiplication) *)
  and v2 = (valeur e2)
  in v1*v2

```

D'une manière générale, il y aura autant de fonctions à coder que de catégories de jugements, et chaque fonction comportera autant de cas qu'il y a de règles sémantiques ou d'axiomes pour la catégorie de jugements correspondante.

1.5 Exercices

Exercice 1.1

On considère un langage qui autorise l'écriture des constantes entières en binaire en les préfixant par le caractère #. Par exemple, on pourra écrire #10010110 pour la constante 150. La syntaxe de ces constantes en binaire est donnée par la grammaire

$$\begin{aligned}
 N &\rightarrow \#B \\
 B &\rightarrow C \mid BC \\
 C &\rightarrow 0 \mid 1
 \end{aligned}$$

Pour déterminer la valeur de ces constantes, on définit au préalable la syntaxe abstraite suivante :

$$\begin{aligned}
 B &\rightarrow \text{Singl}(C) \mid \text{Seq}(B, C) \\
 C &\rightarrow 0 \mid 1
 \end{aligned}$$

Le nombre #1101 sera ainsi représenté par $\text{Seq}(\text{Seq}(\text{Seq}(\text{Singl}(1), 1), 0), 1)$.

1. Associer à chaque règle de grammaire un calcul de syntaxe abstraite.

Correction :

$$\begin{aligned}
 N &\rightarrow \#B && \{B_1\} \\
 B &\rightarrow C && \{\text{Singl}(C_1)\} \\
 B &\rightarrow BC && \{\text{Seq}(B_1, C_1)\} \\
 C &\rightarrow 0 && \{0\} \\
 C &\rightarrow 1 && \{1\}
 \end{aligned}$$

2. Déterminer des règles sémantiques pour l'évaluation des constantes binaires.

Correction :

$$\begin{aligned}
 &\overline{0 \Rightarrow 0} \\
 &\overline{1 \Rightarrow 1} \\
 &\overline{\text{Singl}(c) \Rightarrow c} \\
 &\overline{b \Rightarrow n} \\
 &\overline{\text{Seq}(b, c) \Rightarrow 2n + c}
 \end{aligned}$$

3. Construire l'arbre de dérivation du jugement $\text{Seq}(\text{Seq}(\text{Seq}(\text{Singl}(1), 1), 0), 1) \Rightarrow 13$

Exercice 1.2

On veut généraliser l'exercice précédent de manière à pouvoir écrire dans une base quelconque (entre 2 et 16). On veut écrire les nombres en base b sous la forme $\$b\#n$ où n est une suite de chiffres entre 0 et $b - 1$. Par exemple, on pourra écrire $\$5\#1100$ pour la constante 150 en base 5. La syntaxe de ces constantes est donnée par la pseudo-grammaire

$$\begin{aligned} N &\rightarrow \$ Nat \# D \\ D &\rightarrow C \mid DC \\ C &\rightarrow 0 \mid \dots \mid 9 \mid A \mid \dots \mid F \\ Nat &\rightarrow (0 \mid \dots \mid 9)^+ \end{aligned}$$

1. Définir une syntaxe abstraite pour ce langage.

Correction :

$$\begin{aligned} N &\rightarrow Nombre(Int, D) \\ D &\rightarrow Singl(Char) \mid Seq(D, Char) \end{aligned}$$

2. Associer à chaque règle de grammaire un calcul de syntaxe abstraite.

Correction :

$$\begin{aligned} N &\rightarrow \$Int\#D \quad \{Nombre(val(Int_1), D_1)\} \\ D &\rightarrow C \quad \{Singl(C_1)\} \\ D &\rightarrow DC \quad \{Seq(D_1, C_1)\} \\ C &\rightarrow Char \quad \{val(Char_1)\} \end{aligned}$$

3. Déterminer des règles sémantiques pour l'évaluation des constantes en base quelconque. On utilisera pour les arbres de sorte D un jugement de la forme $b \vdash d \Rightarrow v$ où b est la base, d la syntaxe abstraite représentant la séquence de chiffres et v la valeur. On fera attention à vérifier que la base est entre 2 et 16 et que tous les chiffres sont entre 0 et $b - 1$.

Correction :

$$\frac{}{n \Rightarrow n} \text{ pour tout } n \in \{0, \dots, 9\}$$

$$\frac{}{A \Rightarrow 10}$$

$$\frac{}{B \Rightarrow 11}$$

$$\vdots$$

$$\frac{}{F \Rightarrow 15}$$

$$\frac{c \Rightarrow n}{b \vdash Singl(c) \Rightarrow n} \text{ si } n < b$$

$$\frac{b \vdash d \Rightarrow v \quad c \Rightarrow n}{b \vdash Seq(d, c) \Rightarrow bv + n} \text{ si } n < b$$

$$\frac{b \vdash d \Rightarrow v}{Nombre(b, d) \Rightarrow v} \text{ si } 2 \leq b \leq 16$$

4. Construire l'arbre de dérivation du jugement $a \Rightarrow v$ où a est l'arbre de syntaxe abstraite représentant $\$5\#1432$ et où v est la valeur appropriée, que l'on déterminera justement en construisant cet arbre.

Correction :

$$Nombre(5, Seq(Seq(Seq(Singl(1), 4), 3), 2) \Rightarrow ?$$

5. Donner deux exemples d'expressions concrètes et leur représentations abstraites déclenchant respectivement les deux erreurs de sémantiques possibles. Dans chacun des cas, écrire les portions d'arbres de dérivations en montrant à quel endroit l'erreur se produit.
6. Écrire un interpréteur de ce langage en CAML à l'aide de CAMLYACC.

Chapitre 2

Sémantique des expressions arithmétiques et logiques

Nous allons maintenant étudier un langage simple, que l'on trouve dans tous les langages de programmation, celui des expressions arithmétiques et logiques. Nous nous contenterons tout d'abord de traiter les expressions arithmétiques entières. Nous traiterons ensuite les expressions logiques, avant de considérer le mélange d'expressions numériques et logiques. Le traitement d'expressions arithmétiques quelconques, comportant à la fois des entiers et des réels est beaucoup plus complexe, et nous l'aborderons seulement en fin de chapitre, en ne traitant qu'un sous-cas simple. Dans tous les cas, nous réserverons le traitement des identificateurs au chapitre suivant.

Dans chaque section, nous introduirons d'abord la syntaxe abstraite choisie, avant de considérer les problèmes de typage s'il y a lieu, puis les problèmes d'évaluation.

2.1 Expressions arithmétiques

Nous conservons la syntaxe concrète des chapitres précédents, en nous limitant à l'addition et la multiplication. Ce paragraphe ne fait que regrouper éléments épars, en respectant la démarche :

1. définition de la syntaxe abstraite ;
2. définition de la grammaire concrète non-ambiguë et calcul de la syntaxe abstraite ;
3. définition de la sémantique, typage puis évaluation.

2.1.1 Syntaxe abstraite

La syntaxe abstraite que nous choisissons pour les expressions est la suivante :

$$\begin{aligned} E &\rightarrow +(E, E) \mid \times(E, E) \mid \text{Cte}(\text{Nat}) \\ \text{Nat} &\rightarrow 0 \mid 1 \mid 2 \mid \dots \end{aligned}$$

$$\boxed{
\begin{array}{c}
\overline{\text{Cte}(n) \Rightarrow n} \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{+(e_1, e_2) \Rightarrow v_1 + v_2} \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{\times(e_1, e_2) \Rightarrow v_1 \times v_2}
\end{array}
}$$

FIG. 2.1 – Règles d'évaluation des expressions arithmétiques

La grammaire concrète non-ambiguë accompagnée des règles de calcul de la syntaxe abstraite est alors :

$$\begin{array}{lll}
E & \rightarrow & E + T \quad \{+(E_1, T_1)\} \\
E & \rightarrow & T \quad \{T_1\} \\
T & \rightarrow & T \times F \quad \{\times(T_1, F_1)\} \\
T & \rightarrow & F \quad \{F_1\} \\
F & \rightarrow & \text{Nat} \quad \{\text{Cte}(\text{val}(\text{Nat}_1))\} \\
F & \rightarrow & (E) \quad \{E_1\}
\end{array}$$

2.1.2 Évaluation

Nous pouvons maintenant définir la sémantique des expressions arithmétiques comme une fonction des arbres de syntaxe abstraite dans le domaine sémantique. Les règles d'évaluation sont représentées sur la figure 2.1. Dans ces règles, les opérateurs $+$ et $\times$ apparaissant à droite du symbole $\Rightarrow$ sont les opérations arithmétiques habituelles.

Par exemple, l'évaluation de l'expression arithmétique $(1 + 2) * 3$ est obtenue par l'arbre de dérivation suivant :

$$\frac{\frac{\overline{\text{Cte}(1) \Rightarrow 1} \quad \overline{\text{Cte}(2) \Rightarrow 2}}{+(\text{Cte}(1), \text{Cte}(2)) \Rightarrow 3} \quad \overline{\text{Cte}(3) \Rightarrow 3}}{\times(+(\text{Cte}(1), \text{Cte}(2)), \text{Cte}(3)) \Rightarrow 9}$$

2.2 Expressions logiques

2.2.1 Syntaxe abstraite

La syntaxe abstraite que nous choisissons pour les expressions logiques est la suivante :

$$L \rightarrow \text{And}(L, L) \mid \text{Or}(L, L) \mid \text{Not}(L) \mid \text{True} \mid \text{False}$$

Pour définir une grammaire concrète non-ambiguë pour de telles expressions, il nous faut choisir une priorité entre les opérateurs `and`, `or` et `not`. Le choix usuel dans les langages de programmation est de prendre `not` de priorité supérieure à `and`, lui-même prioritaire sur `or`, et ces deux derniers doivent associer à gauche. La grammaire accompagnée des

$\overline{\text{True} \Rightarrow T}$
$\overline{\text{False} \Rightarrow F}$
$\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{\text{And}(e_1, e_2) \Rightarrow v_1 \wedge v_2}$
$\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{\text{Or}(e_1, e_2) \Rightarrow v_1 \vee v_2}$
$\frac{e \Rightarrow v}{\text{Not}(e) \Rightarrow \neg v}$

FIG. 2.2 – Règles d'évaluation des expressions logiques.

règles de calcul de la syntaxe abstraite est alors :

$L \rightarrow L \text{ or } M$	$\{\text{Or}(L_1, M_1)\}$
$L \rightarrow M$	$\{M_1\}$
$M \rightarrow M \text{ and } N$	$\{\text{And}(M_1, N_1)\}$
$M \rightarrow N$	$\{N_1\}$
$N \rightarrow \text{not } N$	$\{\text{Not}(N_1)\}$
$N \rightarrow O$	$\{O_1\}$
$O \rightarrow \text{true}$	$\{\text{True}\}$
$O \rightarrow \text{false}$	$\{\text{False}\}$
$O \rightarrow (L)$	$\{L_1\}$

2.2.2 Évaluation

Les règles d'évaluation sont représentées sur la figure 2.2. Dans ces règles, les opérateurs $\neg$, $\wedge$ et $\vee$ apparaissant à droite du symbole $\Rightarrow$ sont les opérations booléennes sur l'ensemble $\mathbb{B} = \{T, F\}$ des valeurs booléennes, correspondant respectivement à la négation, la conjonction (et) et la disjonction (ou).

L'évaluation de l'expression logique « `(true and false) or not true` » est obtenue par la dérivation suivante :

$$\frac{\frac{\overline{\text{False} \Rightarrow F} \quad \overline{\text{True} \Rightarrow T}}{\text{And}(\text{True}, \text{False}) \Rightarrow F} \quad \frac{\overline{\text{True} \Rightarrow T}}{\text{Not}(\text{True}) \Rightarrow F}}{\text{Or}(\text{And}(\text{True}, \text{False}), \text{Not}(\text{True})) \Rightarrow F}$$

2.3 Expressions arithmético-logiques

Le mélange des expressions arithmétiques et logiques se produit lorsque l'on fait intervenir des opérateurs de comparaison sur les nombres. Pour simplifier, nous ne considérons ici que les opérateurs $=$ et $>$. Le mélange se produit également en présence d'une opération conditionnelle de la forme « `if then else` » rendant un entier comme résultat.

Nous allons successivement voir deux possibilités de définir la syntaxe abstraite des expressions obtenues, appelées expressions arithmético-logiques. La première est une syntaxe

abstraite à deux sortes, les expressions de sorte « arithmétique » et celles de sorte « booléenne », obtenue grosso modo en juxtaposant la syntaxe abstraite des expressions arithmétiques avec celle des expressions booléennes. Cette approche réduit le typage à sa plus simple expression, types et sortes étant identifiés. La seconde est une syntaxe abstraite à une sorte, et il faudra alors donner des règles de typage. Cette seconde approche est plus robuste, (et elle s'avérera indispensable lorsque le fragment de langage considéré comportera des déclarations d'identificateurs) mais elle nécessite une phase de vérification de types, pour rejeter les expressions sans aucun sens comme « 1 or true » ou « false+2 ». Nous souhaitons en effet qu'avant d'évaluer une expression, toutes les vérifications de types soient effectuées pour garantir l'aboutissement des calculs. Un langage de programmation où le typage garantit que l'évaluation pourra se faire s'appelle un langage *fortement typé*.

2.3.1 Syntaxe abstraite à deux sortes

La syntaxe abstraite à deux sortes est donnée par la grammaire :

$$\begin{aligned}
 L &\rightarrow \text{And}(L, L) \mid \text{Or}(L, L) \mid \text{Not}(L) \mid \text{True} \mid \text{False} \mid = (E, E) \mid > (E, E) \\
 E &\rightarrow \text{If}(L, E, E) \mid + (E, E) \mid \times (E, E) \mid \text{Cte}(\text{Nat}) \\
 \text{Nat} &\rightarrow 0 \mid 1 \mid 2 \mid \dots
 \end{aligned}$$

La grammaire concrète non-ambiguë accompagnée des règles de calcul de la syntaxe abstraite est alors :

$$\begin{array}{ll}
 L \rightarrow L \text{ or } M & \{\text{Or}(L_1, M_1)\} \\
 L \rightarrow M & \{M_1\} \\
 M \rightarrow M \text{ and } N & \{\text{And}(M_1, N_1)\} \\
 M \rightarrow N & \{N_1\} \\
 N \rightarrow \text{not } N & \{\text{Not}(N_1)\} \\
 N \rightarrow O & \{O_1\} \\
 O \rightarrow \text{true} & \{\text{True}\} \\
 O \rightarrow \text{false} & \{\text{False}\} \\
 O \rightarrow (L) & \{L_1\} \\
 O \rightarrow E = E & \{= (E_1, E_2)\} \\
 O \rightarrow E > E & \{> (E_1, E_2)\} \\
 E \rightarrow \text{if } L \text{ then } E \text{ else } E & \{\text{If}(L_1, E_1, E_2)\} \\
 E \rightarrow F & \{F_1\} \\
 F \rightarrow F + G & \{+(F_1, G_1)\} \\
 F \rightarrow G & \{G_1\} \\
 G \rightarrow G * H & \{\times(G_1, H_1)\} \\
 G \rightarrow H & \{H_1\} \\
 H \rightarrow \text{Nat} & \{\text{Cte}(\text{val}(\text{Nat}_1))\} \\
 H \rightarrow (E) & \{E_1\}
 \end{array}$$

Notons que l'ambiguïté d'une expression de la forme « if E then E else E + E » a été résolue en donnant priorité à l'addition.

2.3.2 Évaluation

Les règles d'évaluation sont représentées sur la figure 2.3. Remarquons que les règles d'évaluation des opérations de comparaison sont « conditionnelles » : comme mentionné au paragraphe 1.2 ce sont des schémas de règles qui dénotent l'ensemble des règles obtenues en remplaçant les variables e_1 et e_2 par des expressions abstraites quelconques, et v_1 et v_2 par des entiers satisfaisant la condition donnée.

$$\begin{array}{c}
\overline{\text{True} \Rightarrow T} \\
\\
\overline{\text{False} \Rightarrow F} \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{\text{And}(e_1, e_2) \Rightarrow v_1 \wedge v_2} \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{\text{Or}(e_1, e_2) \Rightarrow v_1 \vee v_2} \\
\\
\frac{e \Rightarrow v}{\text{Not}(e) \Rightarrow \neg v} \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{= (e_1, e_2) \Rightarrow T} \quad \text{si } v_1 = v_2 \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{= (e_1, e_2) \Rightarrow F} \quad \text{si } v_1 \neq v_2 \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{> (e_1, e_2) \Rightarrow T} \quad \text{si } v_1 > v_2 \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{> (e_1, e_2) \Rightarrow F} \quad \text{si } v_1 \not> v_2 \\
\\
\frac{b \Rightarrow T \quad e_1 \Rightarrow v_1}{\text{If}(b, e_1, e_2) \Rightarrow v_1} \\
\\
\frac{b \Rightarrow F \quad e_2 \Rightarrow v_2}{\text{If}(b, e_1, e_2) \Rightarrow v_2} \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{+(e_1, e_2) \Rightarrow v_1 + v_2} \\
\\
\frac{e_1 \Rightarrow v_1 \quad e_2 \Rightarrow v_2}{\times(e_1, e_2) \Rightarrow v_1 \times v_2} \\
\\
\overline{\text{Cte}(n) \Rightarrow n}
\end{array}$$

FIG. 2.3 – Règles d'évaluation des expressions arithmético-logiques.

On obtient l'évaluation de l'expression arithmétique « $2 > 3$ or $2+2=4$ » par la dérivation suivante :

$$\frac{\frac{\overline{\text{Cte}(2) \Rightarrow 2} \quad \overline{\text{Cte}(3) \Rightarrow 3}}{> (\text{Cte}(2), \text{Cte}(3)) \Rightarrow \text{False}} \quad \frac{\frac{\overline{\text{Cte}(2) \Rightarrow 2} \quad \overline{\text{Cte}(2) \Rightarrow 2}}{+ (\text{Cte}(2), \text{Cte}(2)) \Rightarrow 4} \quad \overline{\text{Cte}(4) \Rightarrow 4}}{= (+ (\text{Cte}(2), \text{Cte}(2)), \text{Cte}(4)) \Rightarrow \text{True}}}{\text{Or}(> (\text{Cte}(2), \text{Cte}(3)), = (+ (\text{Cte}(2), \text{Cte}(2)), \text{Cte}(4))) \Rightarrow \text{True}}$$

2.3.3 Syntaxe abstraite à une sorte

Dans les langages de programmation classiques, la grammaire ne fait pas de distinction entre les expressions numériques et booléennes. Il n'y a donc qu'une seule sorte d'expressions abstraites, dont la syntaxe est alors la suivante :

$$\begin{aligned} E &\rightarrow \text{And}(E, E) \mid \text{Or}(E, E) \mid \text{Not}(E) \mid \text{True} \mid \text{False} \\ &\mid \text{If}(E, E, E) \mid = (E, E) \mid > (E, E) \\ &\mid + (E, E) \mid \times (E, E) \mid \text{Cte}(\text{Nat}) \\ \text{Nat} &\rightarrow (0 \mid \dots \mid 9)^+ \end{aligned}$$

Pour écrire la grammaire concrète non-ambiguë, nous devons faire des conventions supplémentaires : on suppose classiquement que les opérateurs de comparaison sont prioritaires sur les opérateurs logiques, mais moins prioritaires que les opérations arithmétiques. Ces conventions permettent de parenthéser correctement une expression comme « $1 + 2 > 4$ and $5 = 6$ » en « $((1 + 2) > 4)$ and $(5 = 6)$ ». La grammaire accompagnée des règles de calcul de la syntaxe abstraite est alors la suivante :

$$\begin{aligned} E &\rightarrow \text{if } E \text{ then } E \text{ else } E & \{\text{If}(E_1, E_2, E_3)\} \\ E &\rightarrow F & \{F_1\} \\ F &\rightarrow F \text{ or } G & \{\text{Or}(F_1, G_1)\} \\ F &\rightarrow G & \{G_1\} \\ G &\rightarrow G \text{ and } H & \{\text{And}(G_1, H_1)\} \\ G &\rightarrow H & \{H_1\} \\ H &\rightarrow \text{not } H & \{\text{Not}(H_1)\} \\ H &\rightarrow I & \{I_1\} \\ I &\rightarrow I = J & \{=(I_1, J_1)\} \\ I &\rightarrow I > J & \{>(I_1, J_1)\} \\ I &\rightarrow J & \{J_1\} \\ J &\rightarrow J + K & \{+(J_1, K_1)\} \\ J &\rightarrow K & \{K_1\} \\ K &\rightarrow K * L & \{\times(K_1, L_1)\} \\ K &\rightarrow L & \{L_1\} \\ L &\rightarrow \text{true} & \{\text{True}\} \\ L &\rightarrow \text{false} & \{\text{False}\} \\ L &\rightarrow \text{Nat} & \{\text{Cte}(\text{val}(\text{Nat}_1))\} \\ L &\rightarrow (E) & \{E_1\} \end{aligned}$$

2.3.4 Typage

Notre grammaire des expressions arithmétiques et logiques admet une seule sorte d'expressions abstraites, alors qu'il y a des expressions de deux types différents, Int et Bool.

$\overline{\text{True} : \text{Bool}}$
$\overline{\text{False} : \text{Bool}}$
$\frac{e_1 : \text{Bool} \quad e_2 : \text{Bool}}{\text{Or}(e_1, e_2) : \text{Bool}}$
$\frac{e_1 : \text{Bool} \quad e_2 : \text{Bool}}{\text{And}(e_1, e_2) : \text{Bool}}$
$\frac{e : \text{Bool}}{\text{Not}(e) : \text{Bool}}$
$\frac{e_1 : \text{Int} \quad e_2 : \text{Int}}{= (e_1, e_2) : \text{Bool}}$
$\frac{e_1 : \text{Int} \quad e_2 : \text{Int}}{> (e_1, e_2) : \text{Bool}}$
$\frac{b : \text{Bool} \quad e_1 : \text{Int} \quad e_2 : \text{Int}}{\text{If}(b, e_1, e_2) : \text{Int}}$
$\frac{e_1 : \text{Int} \quad e_2 : \text{Int}}{+ (e_1, e_2) : \text{Int}}$
$\frac{e_1 : \text{Int} \quad e_2 : \text{Int}}{\times (e_1, e_2) : \text{Int}}$
$\overline{\text{Cte}(n) : \text{Int}}$

FIG. 2.4 – *Typage des expressions arithmético-logiques*

Le *type* d'une expression sera alors déterminé au cours de la phase d'analyse sémantique, définie au moyen des règles de typage de la figure 2.4.

Le jugement de typage de l'expression arithmético-logique « $2 > 3$ or $2+2=4$ » est alors le suivant :

$$\frac{\frac{\overline{\text{Cte}(2) : \text{Int}} \quad \overline{\text{Cte}(3) : \text{Int}}}{> (\text{Cte}(2), \text{Cte}(3)) : \text{Bool}} \quad \frac{\frac{\overline{\text{Cte}(2) : \text{Int}} \quad \overline{\text{Cte}(2) : \text{Int}}}{+ (\text{Cte}(2), \text{Cte}(2)) : \text{Int}} \quad \overline{\text{Cte}(4) : \text{Int}}}{= (+ (\text{Cte}(2), \text{Cte}(2)), \text{Cte}(4)) : \text{Bool}}}{\text{Or}(> (\text{Cte}(2), \text{Cte}(3)), = (+ (\text{Cte}(2), \text{Cte}(2)), \text{Cte}(4))) : \text{Bool}}$$

Sur une expression mal typée comme par exemple « $(2=3)+4$ », l'échec du typage et mis en évidence par l'impossibilité de compléter l'arbre de dérivation :

$$\frac{\frac{\overline{\text{Cte}(2) : \text{Int}} \quad \overline{\text{Cte}(3) : \text{Int}}}{= (\text{Cte}(2), \text{Cte}(3)) : \text{Bool}} \quad \overline{\text{Cte}(4) : \text{Int}}}{+ (= (\text{Cte}(2), \text{Cte}(3)), \text{Cte}(4)) : ?}$$

2.3.5 Évaluation

La syntaxe à une sorte possède plus d'expressions possibles que la syntaxe à deux sortes, mais le typage permet de discriminer les expressions qui n'ont pas de sens (mal typées) de celles qui en ont un, qui correspondent alors exactement avec les expressions de la grammaire à deux sortes. Il n'y a donc aucune différence avec l'évaluation des expressions arithmético-logiques décrite au paragraphe 2.3.2, puisque les expressions de syntaxe abstraite correctement typées engendrées par la grammaire à une sorte, sont identiques à celles engendrées par la grammaire à deux sortes.

2.3.6 Notion de typage fort

Dans la section sur les expressions arithmético-logiques, nous avons du faire face à deux situations différentes. Dans le cas des expressions à deux sortes, toute expression de syntaxe abstraite était évaluable, il n'y avait donc nul besoin de contrôle de type. Au contraire, dans le cas des expressions à une sorte, certaines expressions de syntaxe abstraite ne pouvaient pas être évaluées, comme l'addition des expressions $\text{Cte}(n)$ et True . Les règles de typage servent précisément à éliminer toutes les expressions de syntaxe abstraite pour lesquelles il n'existe pas de dérivation d'évaluation.

Définition 2.1 *Un langage de programmation défini par sa syntaxe abstraite, ses règles de typage et ses règles d'évaluation sera dit fortement typé, s'il existe une dérivation d'évaluation pour toute expression pour laquelle il existe une dérivation de typage.*

La définition qui précède s'applique parfaitement au langage des expressions arithmétiques de cette section, dont on peut montrer sans difficulté qu'il est fortement typé:

Théorème 2.2 *Le langage des expressions arithmético-logiques à une sorte est fortement typé.*

Notre définition du typage fort est toutefois un peu simpliste, car elle ne prend pas en compte la possibilité qu'une exécution soit interrompue par la *levée d'une exception*, comme la division par zéro dans un langage d'expressions arithmétiques avec division. De façon générale, la notion d'exception est très directement liée celle de fonction partielle. Ainsi, la division est une fonction partielle, car elle n'est pas définie pour toute valeur de son second

argument. De même l'accès aux éléments d'un tableau, ou certaines fonctions de traitement des chaînes de caractères sont des fonctions partielles. Formaliser les notions de fonction partielle et d'exception sort du cadre de ce cours. Mais il faut retenir que la notion de typage fort doit prendre en compte l'existence de fonctions partielles, la levée d'une exception étant considérée comme l'aboutissement d'une dérivation d'évaluation.

Un langage de programmation fortement typé garantit que tout programme qui se compile sans erreur s'exécutera jusqu'à son terme (sauf à lever une exception, comme on vient de le voir). À titre d'exemple, le langage PASCAL n'est pas fortement typé, à cause des types pointeurs d'une part et des enregistrements avec variantes d'autre part. Le langage C ne l'est pas non plus, à cause des pointeurs, des types union, et aussi de l'absence de vérification des types des fonctions d'un module à l'autre. Les langages fortement typés les plus connus sont ceux de la famille ML (en particulier, le langage CAML), ADA et JAVA.

Notons pour terminer cette discussion que la notion de typage fort ne garantit pas que le programme soit conforme à sa spécification formelle lorsqu'elle existe, ou à une spécification informelle décrite dans un cahier des charges. Elle ne fait que garantir que l'exécution du programme *à l'aide des règles sémantiques* peut être menée à son terme. Elle ne garantit pas non plus que l'exécution du programme sur une machine particulière ira jusqu'à son terme. Il se peut en effet que les ressources de la machine (pile d'exécution, par exemple) soient insuffisantes pour une exécution donnée.

2.4 Expressions arithmétiques surchargées

La grammaire non ambiguë des expressions arithmétiques que nous avons utilisée admet un seul type d'expressions arithmétiques, les expressions arithmétiques entières, reconnues par un arbre de racine E . Mais la plupart des langages réels, comme PASCAL, admettent ce que l'on appelle de la *surcharge d'opérateurs*, c'est-à-dire que l'argument d'une opération arithmétique comme $+$ pourra être indifféremment un entier ou un réel (représenté par exemple en virgule flottante simple précision). Nous devons donc gérer la coexistence de plusieurs types d'expressions arithmétiques, les expressions arithmétiques entières, les expressions arithmétiques réelles simple précision, et, peut-être, les expressions arithmétiques réelles double précision, voire les expressions arithmétiques entières en précision arbitraire. Les constantes arithmétiques de ces différents types pourront être aisément distinguées par l'analyseur lexical : les constantes entières sont des suites de chiffres, comme 31415926535, alors que les constantes réelles simple précision sont constituées de deux suites de chiffres séparées par un point (dit décimal), comme 3.1415926535, suivies éventuellement d'un exposant, comme 6.02e23. Nous supposons par la suite que l'analyseur lexical renvoie le token *Nat* pour une constante entière, et le token *Real* pour une constante réelle.

2.4.1 Syntaxe abstraite

La grammaire abstraite des expressions arithmétiques surchargées est la suivante :

$$\begin{aligned} E &\rightarrow +(E, E) \mid \times (E, E) \mid \text{Cte}(\text{Nat}) \mid \text{Ctreal}(\text{Real}) \\ \text{Nat} &\rightarrow (0 \mid \dots \mid 9)^+ \\ \text{Real} &\rightarrow \text{cf. exercice ??} \end{aligned}$$

La grammaire concrète non-ambiguë accompagnée des règles de calcul de la syntaxe abstraite est alors :

$$\begin{array}{lll}
 E & \rightarrow & E + T \quad \{+(E_1, T_1)\} \\
 E & \rightarrow & T \quad \{T_1\} \\
 T & \rightarrow & T \times F \quad \{\times(T_1, F_1)\} \\
 T & \rightarrow & F \quad \{F_1\} \\
 F & \rightarrow & Nat \quad \{\text{Cte}(\text{val}(Nat_1))\} \\
 F & \rightarrow & Real \quad \{\text{Ctreal}(\text{val}(Real_1))\} \\
 F & \rightarrow & (E) \quad \{E_1\}
 \end{array}$$

2.4.2 Typage

La représentation machine d'un même nombre, par exemple le nombre 2, sera toujours très différente suivant qu'il s'agit d'un 2 entier (dont la représentation binaire sur 32 bits sera en général 000000000000000000000000000010, le bit le plus à gauche donnant le signe, 0 pour positif, et 1 pour négatif), ou d'un 2 réel (dont la représentation binaire sur 32 bits pourra être par exemple 010000000000000000000000000010, le bit le plus à gauche donnant le signe, les 20 bits suivants la mantisse, le bit suivant le signe de l'exposant, et les 10 derniers bits donnant la valeur de l'exposant). Les machines savent additionner des nombres dans une même représentation, par exemple des entiers, mais pas des nombres de deux représentations différentes. La raison n'est pas de principe, il s'agit d'un simple problème de coût : savoir additionner des entiers et des réels nécessiterait deux additionneurs supplémentaires (entier + réel, réel + entier), ce qui coûte cher. En présence de réels double précision, il faudrait cette fois 8 additionneurs au lieu de 3.

Il y a deux solutions à ce problème : soit on interdit le mélange d'entiers et de réels ; soit la responsabilité du choix de l'additionneur adéquat, ainsi que du choix de la représentation la plus adéquate de la donnée est déléguée au logiciel, interprète ou compilateur. Dans le premier cas, une simple analyse de type permettra de décider quel circuit d'addition doit être choisi pour l'évaluation. Le second cas est plus complexe, le principe est d'additionner des entiers chaque fois que cela est possible, et des réels quand on ne peut pas faire autrement. Notre choix, qui est aussi celui de ML, est d'interdire les mélanges d'entiers et de réels¹.

Notons enfin que, même dans le cas simple que nous avons choisi, le résultat du typage n'est pas une simple réponse par oui ou non, puisqu'il faudra distinguer les opérations de même nom, comme + et $\times$, suivant le type de leurs arguments. La phase de typage va donc devoir calculer une nouvelle syntaxe abstraite, dans laquelle figurent des informations de type sous la forme d'opérations d'addition et de multiplication désambiguïsées, c'est-à-dire indexées par le type de leurs arguments. C'est cette syntaxe abstraite là, décrite par la grammaire ci-dessous, qui sera utilisée à l'évaluation.

$$\begin{array}{l}
 E \rightarrow +_{int}(E, E) \mid +_{real}(E, E) \mid \times_{int}(E, E) \mid \times_{real}(E, E) \\
 \mid \text{Cte}(Nat) \mid \text{Ctreal}(Real)
 \end{array}$$

Les jugements de typage donnés figure 2.5 sont donc un peu plus complexes, car le domaine de valeur sera $D^v = \{\text{Int}, \text{Real}\} \times A$ où A est l'ensemble des arbres abstraits de la grammaire ci-dessus. Nous les noterons sous la forme $e : t \Rightarrow e'$, où e est l'expression de syntaxe abstraite, t est le type de e , et e' est l'expression de syntaxe abstraite calculée.

Le jugement de typage de l'expression arithmétique $(1 + 3) \times 3.14$ est représenté à la figure 2.6. On remarque que ce jugement échoue, à cause du mélange de types.

1. Nous parlons ici de ML standard, CAML-light, lui, interdit complètement la surcharge.

$$\begin{array}{c}
\frac{e_1 : \text{Int} \Rightarrow e'_1 \quad e_2 : \text{Int} \Rightarrow e'_2}{+(e_1, e_2) : \text{Int} \Rightarrow +_{int}(e'_1, e'_2)} \\
\\
\frac{e_1 : \text{Real} \Rightarrow e'_1 \quad e_2 : \text{Real} \Rightarrow e'_2}{+(e_1, e_2) : \text{Real} \Rightarrow +_{real}(e'_1, e'_2)} \\
\\
\frac{e_1 : \text{Int} \Rightarrow e'_1 \quad e_2 : \text{Int} \Rightarrow e'_2}{\times(e_1, e_2) : \text{Int} \Rightarrow \times_{int}(e'_1, e'_2)} \\
\\
\frac{e_1 : \text{Real} \Rightarrow e'_1 \quad e_2 : \text{Real} \Rightarrow e'_2}{\times(e_1, e_2) : \text{Real} \Rightarrow \times_{real}(e'_1, e'_2)} \\
\\
\frac{}{\overline{\text{Cte}(n) : \text{Int} \Rightarrow \text{Cte}(n)}} \\
\\
\frac{}{\overline{\text{Ctereal}(r) : \text{Real} \Rightarrow \text{Ctereal}(r)}}
\end{array}$$

FIG. 2.5 – *Typage des expressions arithmétiques surchargées*

$$\frac{\frac{\overline{\text{Cte}(3) : \text{Int} \Rightarrow \text{Cte}(3)} \quad \overline{\text{Cte}(1) : \text{Int} \Rightarrow \text{Cte}(1)}}{+(\text{Cte}(1), \text{Cte}(3)) : \text{Int} \Rightarrow +_{int}(\text{Cte}(1), \text{Cte}(3))} \quad \frac{}{\overline{\text{Ctereal}(3.14) : \text{Real} \Rightarrow \text{Ctereal}(3.14)}}}{\times(+(\text{Cte}(1), \text{Cte}(3)), \text{Ctereal}(3.14)) : ? \Rightarrow ?}$$

FIG. 2.6 – *Échec du jugement de typage de $(1 + 3) \times 3.14$*

2.4.3 Évaluation

Le typage d'une expression est toujours *statique*: il ne dépend pas de l'évaluation de l'expression. Par contre, l'évaluation des expressions peut dépendre du typage par le biais du calcul d'une nouvelle syntaxe abstraite comme précédemment. Les phases de typage et d'évaluation restent toutefois clairement séparées.

L'évaluation des expressions de syntaxe abstraite (décorée par les types) d'une arithmétique surchargée n'est pas très différente de l'évaluation déjà donnée pour les expressions non-surchargées. Nous laissons au lecteur le soin d'en élaborer les règles.

2.5 Exercices

Exercice 2.1

On prolonge l'exercice 5.1.

1. Coder les règles d'évaluation des expressions booléennes en une fonction qui étant donnée une expression booléenne abstraite retourne sa valeur.
2. Écrire un programme CAML qui lit au clavier une expression booléenne et affiche sa valeur à l'écran.

Exercice 2.2

On prolonge l'exercice 5.2.

1. Coder les règles d'évaluation des expressions arithmétiques en une fonction qui étant donnée une expression arithmétique abstraite retourne sa valeur.
2. Écrire un programme CAML qui lit au clavier une expression arithmétique et affiche sa valeur à l'écran.

Exercice 2.3

Reprendre et combiner les deux exercices précédents pour écrire un programme CAML qui lit au clavier une expression arithmético-logique et affiche sa valeur à l'écran.

Exercice 2.4

Définir les syntaxe concrète puis abstraite, les règles de typage et d'évaluation pour un langage d'expressions mélangeant les entiers et les chaînes de caractères, muni d'un opérateur de concaténation noté « $\cdot$ », d'une fonction `len` retournant la longueur d'une chaîne, et d'une fonction `sub` à trois arguments telles que `sub( $s, n, m$ )` retourne la sous-chaîne de s entre les n -ième et m -ième caractères. Donner les dérivations des jugements de typage et d'évaluation de

```
sub("bonjour", len("bon")+1, 7) . "soir"
```

Exercice 2.5

1. Montrer que le langage des expressions arithmétiques surchargées est fortement typé.
2. Définir un langage d'expressions arithmético-logiques avec surcharge des opérateurs arithmétiques. On donnera d'abord, (i) sa syntaxe abstraite à une sorte, (ii) sa syntaxe concrète accompagnée des règles de calcul de la syntaxe abstraite, (iii) ses règles de typage, (iv) ses règles d'évaluation. On montrera pour terminer que le langage obtenu est fortement typé.

Chapitre 3

Variables et environnements

Ce chapitre s'intéresse aux problèmes de sémantiques liés à l'utilisation de variables.

Nous avons supposé jusqu'à maintenant que nos expressions arithmétiques ne contenaient pas d'identificateurs. Nous allons maintenant définir des *expressions arithmétiques étendues*, en introduisant dans la grammaire une construction **let**. Nous supposerons dans un premier temps comme dans le paragraphe 2.1 que nous avons un seul type de constantes, les constantes entières. Les jugements de typage que nous allons introduire consisteront simplement à détecter les variables non déclarées. Dans un second temps, nous réintroduirons des booléens dans les expressions, et les jugements de typage détecteront également les variables dont le type déclaré n'est pas le bon.

3.1 Cas des expressions arithmétiques seules

Pour définir des variables locales à une expression, nous introduisons la syntaxe classique de ML : $\text{let } x = \text{expr}_1 \text{ in } \text{expr}_2$. La syntaxe concrète est donc

$$\begin{aligned} E &\rightarrow \text{let } Id = E \text{ in } E \mid E + E \mid E \times E \mid Nat \mid Id \mid (E) \\ Id &\rightarrow \text{cf. exercice ??} \end{aligned}$$

où *Id* reconnaît les identificateurs. Le **let** introduit une nouvelle ambiguïté dans cette grammaire : une expression comme « **let** $x = 1$ **in** $x + 2$ » peut être parenthésée comme « **(let** $x = 1$ **in** x) $+ 2$ » ou bien « **let** $x = 1$ **in** $(x + 2)$ ». Le choix que nous faisons, celui qui est classique, est de donner priorité aux opérations par rapport au **let**, et le deuxième parenthésage est donc celui souhaité. Nous commençons par donner une grammaire concrète non-ambiguë en conséquence, après avoir défini au préalable la syntaxe abstraite voulue.

3.1.1 Syntaxe abstraite

La syntaxe abstraite choisie pour les expressions arithmétiques étendues est donnée par la grammaire

$$E \rightarrow \text{Let}(String, E, E) \mid +(E, E) \mid \times(E, E) \mid \text{Cte}(Int) \mid \text{Var}(String)$$

La grammaire concrète non-ambiguë accompagnée des règles de calcul de la syntaxe abstraite est alors :

$E \rightarrow \text{let } Id = E \text{ in } E$	$\{\text{Let}(\text{val}(Id_1), E_1, E_2)\}$
$E \rightarrow F$	$\{F_1\}$
$F \rightarrow F + G$	$\{+(F_1, G_1)\}$
$F \rightarrow G$	$\{G_1\}$
$G \rightarrow G * H$	$\{\times(G_1, H_1)\}$
$G \rightarrow H$	$\{H_1\}$
$H \rightarrow \text{Nat}$	$\{\text{Cte}(\text{val}(\text{Nat}_1))\}$
$H \rightarrow Id$	$\{\text{Var}(\text{val}(Id_1))\}$
$H \rightarrow (E)$	$\{E_1\}$

3.1.2 Typage

On pourrait penser que le typage des expressions arithmétiques étendues est très simple, dans la mesure où le seul type possible est Int . En fait, il est plus complexe qu'il n'y parait, car il se peut qu'une expression ne soit pas typable, lorsqu'un identificateur est utilisé dans une expression arithmétique sans être déclaré auparavant dans un **let**. Si cette vérification n'avait pas lieu, le typage n'assurerait pas son rôle, qui est de garantir que le programme peut être évalué. Cette vérification demande l'utilisation d'un *environnement de liaison* permettant de stocker le type des identificateurs, technique que nous serons également amenés à utiliser pour l'évaluation.

On peut se demander pourquoi la grammaire n'est pas modifiée afin d'éviter cet inconvénient : c'est parce que ce n'est pas possible avec une grammaire hors-contexte (cf. exercice 3.2).

Définition 3.1 *Un environnement de typage est une liste (donc ordonnée) de couples (identificateur, type), notée $\{x_1 : t_1, \dots, x_n : t_n\}$. On notera par $\{\}$ l'environnement vide, et par $\Gamma \cdot \Gamma'$ la concaténation des environnements Γ et Γ' . Le même identificateur peut apparaître plusieurs fois dans un environnement, auquel cas la notation $(x, v) \in \Gamma$ signifie que $\Gamma = \Gamma_1 \cdot \{(x, v)\} \cdot \Gamma_2$ et qu'aucune paire de la forme (x, v') ne figure dans Γ_2 .*

Notons que par conséquent, les environnements se comportent comme des piles vis-à-vis de l'opérateur d'appartenance, où la tête de pile est écrite à droite.

Définition 3.2 *Un jugement de typage est un jugement (Γ, e, t) que l'on écrit $\Gamma \vdash e : t$, où Γ est un environnement de typage (le domaine de contexte est donc l'ensemble des environnements de typage), e est une expression arithmétique abstraite étendue, et t est un type (le domaine de valeurs est donc l'ensemble des types).*

Pour l'instant l'ensemble des types est réduit à $\{\text{Int}\}$. Ce ne sera pas toujours le cas, en particulier s'il y a des expressions arithmético-logiques ou bien des expressions arithmétiques surchargées. Les règles de typage sont décrites à la figure 3.1.

La figure 3.2 montre le typage de l'expression arithmétique étendue **let** $n = 2$ **in** **let** $m = n + 1$ **in** $m + n$ dans l'environnement vide. La dérivation de typage est présentée en trois parties, de manière à tenir dans la largeur de la page. Les références (1) et (2) indiquent où se poursuit la dérivation.

3.1.3 Évaluation

Afin de traiter l'évaluation de ces expressions arithmétiques étendues, nous utilisons un nouveau type de jugement, qui fait appel à un environnement de liaison dont le but est de

$$\boxed{
\begin{array}{c}
\overline{\Gamma \vdash \text{Var}(x) : \text{Int}} \text{ si } x : \text{Int} \in \Gamma \\
\\
\overline{\Gamma \vdash \text{Cte}(n) : \text{Int}} \\
\\
\frac{\Gamma \vdash e_1 : \text{Int} \quad \Gamma \vdash e_2 : \text{Int}}{\Gamma \vdash \times(e_1, e_2) : \text{Int}} \\
\\
\frac{\Gamma \vdash e_1 : \text{Int} \quad \Gamma \vdash e_2 : \text{Int}}{\Gamma \vdash +(e_1, e_2) : \text{Int}} \\
\\
\frac{\Gamma \vdash e_1 : \text{Int} \quad \Gamma \cdot \{x : \text{Int}\} \vdash e_2 : \text{Int}}{\Gamma \vdash \text{Let}(x, e_1, e_2) : \text{Int}}
\end{array}
}$$

FIG. 3.1 – *Typage des expressions arithmétiques étendues*

$$\boxed{
\begin{array}{l}
(1) : \frac{\overline{\{n : \text{Int}\} \vdash \text{Var}(n) : \text{Int}} \quad \overline{\{n : \text{Int}\} \vdash \text{Cte}(1) : \text{Int}}}{\{n : \text{Int}\} \vdash +(\text{Var}(n), \text{Cte}(1)) : \text{Int}} \\
\\
(2) : \frac{\overline{\{n : \text{Int}, m : \text{Int}\} \vdash \text{Var}(m) : \text{Int}} \quad \overline{\{n : \text{Int}, m : \text{Int}\} \vdash \text{Var}(n) : \text{Int}}}{\{n : \text{Int}, m : \text{Int}\} \vdash +(\text{Var}(m), \text{Var}(n)) : \text{Int}} \\
\\
\frac{\overline{\{\} \vdash \text{Cte}(2) : \text{Int}} \quad \frac{(1) \quad (2)}{\{n : \text{Int}\} \vdash \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n))) : \text{Int}}}{\{\} \vdash \text{Let}(n, \text{Cte}(2), \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n)))) : \text{Int}}
\end{array}
}$$

FIG. 3.2 – *Typage de l'expression arithmétique étendue $\text{let } n = 2 \text{ in let } m = n + 1 \text{ in } m + n$*

$$\boxed{
\begin{array}{c}
\frac{}{\Gamma \vdash \text{Var}(x) \Rightarrow v} \text{ si } x = v \in \Gamma \\
\\
\frac{}{\Gamma \vdash \text{Cte}(n) \Rightarrow n} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v \quad \Gamma \vdash e_2 \Rightarrow v'}{\Gamma \vdash \times(e_1, e_2) \Rightarrow v \times v'} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v \quad \Gamma \vdash e_2 \Rightarrow v'}{\Gamma \vdash +(e_1, e_2) \Rightarrow v + v'} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \cdot \{x = v_1\} \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{Let}(x, e_1, e_2) \Rightarrow v_2}
\end{array}
}$$

FIG. 3.3 – Évaluation des expressions arithmétiques étendues

stocker cette fois les valeurs des identificateurs.

Définition 3.3 Un environnement de valeurs est une liste de couples (identificateur, valeur), notée $\{x_1 = v_1, \dots, x_n = v_n\}$.

Un jugement d'évaluation est un jugement (Γ, e, v) où Γ est un environnement de valeurs, e une expression arithmétique abstraite étendue, et v est une valeur.

Étant donnée une expression arithmétique étendue e , le jugement d'évaluation s'écrit $\Gamma \vdash e \Rightarrow v$, où v est une valeur appartenant au domaine d'interprétation, c'est-à-dire les entiers naturels.

Les règles d'évaluation sont décrites sur la figure 3.3. La règle du **let** présente deux particularités. Tout d'abord, la valeur v_1 calculée dans la branche gauche de l'arbre de dérivation est réutilisée dans la branche droite. Cela impose un ordre, gauche-droit, sur le calcul de l'arbre de dérivation. Deuxièmement, l'identificateur x défini par l'expression **let** $x = e_1$ **in** e_2 prend la valeur v_1 dans l'expression e_2 , et dans elle seule. On dit que la *portée* du **let** est l'expression e_2 , ou que les occurrences de x dans e_2 sont dans la portée du **let**. Il se peut que cette expression **let** soit elle-même à l'intérieur d'une autre expression **let**, auquel cas il est important de lier les diverses occurrences de x aux bonnes valeurs. Cela est expliqué en détail plus loin, après les exemples.

Théorème 3.4 Le langage des expressions arithmétiques étendues est fortement typé.

La figure 3.4 montre l'évaluation de l'expression arithmétique étendue **let** $n = 2$ **in** **let** $m = n + 1$ **in** $m + n$ dans l'environnement vide.

Ces exemples montrent que les identificateurs définis au moyen de la construction **let** ont une *portée* bien précise, qui est définie par le parenthésage des expressions **let**. Cette construction permet d'avoir ce que l'on appelle des *variables locales*, fondamentalement différentes des variables globales existant dans certains langages de programmation. Une conséquence importante est que le symbole d'égalité que nous avons utilisé dans la construction (**let** $x = e_1$ **in** e_2) désigne une *définition* ou une *redéfinition* de l'identificateur x , et non pas une *affectation*. En effet, une affectation consiste à écrire une valeur dans la cellule mémoire associée à l'identificateur x : toute valeur écrite antérieurement sera donc remplacée par la nouvelle (on dit *écrasée*). Une redéfinition consiste en une association temporaire d'une valeur à l'identificateur x . Elle ne doit surtout pas effacer une association plus ancienne encore *active*, c'est-à-dire susceptible d'être utilisée ultérieurement. L'exemple de la

$$\begin{array}{l}
(3) : \frac{\frac{\{n = 2, m = 3\} \vdash \text{Var}(m) \Rightarrow 3}{\{n = 2, m = 3\} \vdash +(\text{Var}(m), \text{Var}(n)) \Rightarrow 5} \quad \frac{\{n = 2, m = 3\} \vdash \text{Var}(n) \Rightarrow 2}{\{n = 2, m = 3\} \vdash +(\text{Var}(m), \text{Var}(n)) \Rightarrow 5}}{\{n = 2, m = 3\} \vdash +(\text{Var}(m), \text{Var}(n)) \Rightarrow 5} \\
(2) : \frac{\frac{\{n = 2\} \vdash \text{Var}(n) \Rightarrow 2}{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(1)) \Rightarrow 3} \quad \frac{\{n = 2\} \vdash \text{Cte}(1) \Rightarrow 1}{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(1)) \Rightarrow 3}}{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(1)) \Rightarrow 3} \\
(1) : \frac{\frac{\frac{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(1)) \Rightarrow 3}{\{n = 2\} \vdash \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n))) \Rightarrow 5} \quad \frac{\frac{\{n = 2, m = 3\} \vdash +(\text{Var}(m), \text{Var}(n)) \Rightarrow 5}{\{n = 2, m = 3\} \vdash \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n))) \Rightarrow 5}}{\{n = 2, m = 3\} \vdash \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n))) \Rightarrow 5}}{\{n = 2\} \vdash \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n))) \Rightarrow 5} \\
\frac{\frac{\{\} \vdash \text{Cte}(2) \Rightarrow 2}{\{\} \vdash \text{Let}(n, \text{Cte}(2), \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n)))) \Rightarrow 5} \quad \frac{\frac{\{n = 2\} \vdash \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n))) \Rightarrow 5}{\{\} \vdash \text{Let}(n, \text{Cte}(2), \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n)))) \Rightarrow 5}}{\{\} \vdash \text{Let}(n, \text{Cte}(2), \text{Let}(m, +(\text{Var}(n), \text{Cte}(1)), +(\text{Var}(m), \text{Var}(n)))) \Rightarrow 5}
\end{array}$$

FIG. 3.4 – Évaluation de l'expression arithmétique étendue *let* $n = 2$ *in* *let* $m = n + 1$ *in* $m + n$

$$\begin{array}{l}
(3) : \frac{\frac{\{n = 2\} \vdash \text{Var}(n) \Rightarrow 2}{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(2)) \Rightarrow 4} \quad \frac{\{n = 2\} \vdash \text{Cte}(2) \Rightarrow 2}{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(2)) \Rightarrow 4}}{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(2)) \Rightarrow 4} \\
(2) : \frac{\frac{\frac{\{n = 2\} \vdash +(\text{Var}(n), \text{Cte}(2)) \Rightarrow 4}{\{n = 2\} \vdash \text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)) \Rightarrow 4} \quad \frac{\{n = 2, n = 4\} \vdash \text{Var}(n) \Rightarrow 4}{\{n = 2, n = 4\} \vdash \text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)) \Rightarrow 4}}{\{n = 2\} \vdash \text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)) \Rightarrow 4} \\
(1) : \frac{\frac{\frac{\frac{\{n = 2\} \vdash \text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)) \Rightarrow 4}{\{n = 2\} \vdash +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n)) \Rightarrow 6} \quad \frac{\frac{\{n = 2\} \vdash \text{Var}(n) \Rightarrow 2}{\{n = 2\} \vdash +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n)) \Rightarrow 6}}{\{n = 2\} \vdash +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n)) \Rightarrow 6}}{\{n = 2\} \vdash +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n)) \Rightarrow 6} \\
\frac{\frac{\{\} \vdash \text{Cte}(2) \Rightarrow 2}{\{\} \vdash \text{Let}(n, \text{Cte}(2), +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n))) \Rightarrow 6} \quad \frac{\frac{\{n = 2\} \vdash +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n)) \Rightarrow 6}{\{\} \vdash \text{Let}(n, \text{Cte}(2), +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n))) \Rightarrow 6}}{\{\} \vdash \text{Let}(n, \text{Cte}(2), +(\text{Let}(n, +(\text{Var}(n), \text{Cte}(2)), \text{Var}(n)), \text{Var}(n))) \Rightarrow 6}
\end{array}$$

FIG. 3.5 – Évaluation de l'expression arithmétique étendue *let* $n = 2$ *in* (*let* $n = n + 2$ *in* n) $+ n$

figure 3.5 illustre le problème en imbriquant deux constructions *let* portant sur le même nom d'identificateur.

3.2 Les expressions arithmético-logiques

L'extension de l'étude précédente au cas des expressions arithmético-logiques ne pose pas de difficulté supplémentaires, il s'agit simplement d'introduire le type *Bool* dans l'ensemble des types.

On obtient la grammaire abstraite :

$$\begin{aligned}
 E &\rightarrow \text{And}(E, E) \mid \text{Or}(E, E) \mid \text{Not}(E) \mid \text{True} \mid \text{False} \\
 &\mid \text{If}(E, E, E) \mid \text{Let}(Id, E, E) \mid =(E, E) \mid >(E, E) \\
 &\mid +(E, E) \mid \times(E, E) \mid \text{Cte}(Nat) \mid \text{Var}(Id) \\
 Nat &\rightarrow (0 \mid \dots \mid 9)^+ \\
 Id &\rightarrow (a \mid \dots \mid z)^+
 \end{aligned}$$

La grammaire concrète non ambiguë accompagnée des calculs de syntaxe abstraite devient :

$$\begin{aligned}
 E &\rightarrow \text{if } E \text{ then } E \text{ else } E && \{\text{If}(E_1, E_2, E_1)\} \\
 E &\rightarrow \text{let } Id = E \text{ in } E && \{\text{Let}(\text{val}(Id_1), E_1, E_1)\} \\
 E &\rightarrow F && \{F_1\} \\
 F &\rightarrow F \text{ or } G && \{\text{Or}(F_1, G_1)\} \\
 F &\rightarrow G && \{G_1\} \\
 G &\rightarrow G \text{ and } H && \{\text{And}(G_1, H_1)\} \\
 G &\rightarrow H && \{H_1\} \\
 H &\rightarrow \text{not } H && \{\text{Not}(H_1)\} \\
 H &\rightarrow I && \{I_1\} \\
 I &\rightarrow I = J && \{=(I_1, J_1)\} \\
 I &\rightarrow I > J && \{>(I_1, J_1)\} \\
 I &\rightarrow J && \{J_1\} \\
 J &\rightarrow J + K && \{+(J_1, K_1)\} \\
 J &\rightarrow K && \{K_1\} \\
 K &\rightarrow K * L && \{\times(K_1, L_1)\} \\
 K &\rightarrow L && \{L_1\} \\
 L &\rightarrow \text{true} && \{\text{True}\} \\
 L &\rightarrow \text{false} && \{\text{False}\} \\
 L &\rightarrow Nat && \{\text{Cte}(\text{val}(Nat_1))\} \\
 L &\rightarrow Id && \{\text{Var}(\text{val}(Id_1))\} \\
 L &\rightarrow (E) && \{E_1\}
 \end{aligned}$$

Notons tout de même qu'ici, il ne serait définitivement pas possible d'étendre la grammaire à deux sortes du paragraphe 2.3.1 car le fragment de grammaire

$$\begin{aligned}
 E &\rightarrow Id \\
 L &\rightarrow Id
 \end{aligned}$$

introduirait forcément une ambiguïté.

3.2.1 Typage

Les règles de typage sont données figure 3.6.

3.2.2 Évaluation

Les règles d'évaluation sont présentées figure 3.7.

Théorème 3.5 *Le langage des expressions arithmético-logiques étendues est fortement typé.*

$\overline{\Gamma \vdash \text{True} : \text{Bool}}$
$\overline{\Gamma \vdash \text{False} : \text{Bool}}$
$\frac{\Gamma \vdash e_1 : \text{Bool} \quad \Gamma \vdash e_2 : \text{Bool}}{\Gamma \vdash \text{Or}(e_1, e_2) : \text{Bool}}$
$\frac{\Gamma \vdash e_1 : \text{Bool} \quad \Gamma \vdash e_2 : \text{Bool}}{\Gamma \vdash \text{And}(e_1, e_2) : \text{Bool}}$
$\frac{\Gamma \vdash e : \text{Bool}}{\Gamma \vdash \text{Not}(e) : \text{Bool}}$
$\frac{\Gamma \vdash e_1 : \text{Int} \quad \Gamma \vdash e_2 : \text{Int}}{\Gamma \vdash = (e_1, e_2) : \text{Bool}}$
$\frac{\Gamma \vdash e_1 : \text{Int} \quad \Gamma \vdash e_2 : \text{Int}}{\Gamma \vdash > (e_1, e_2) : \text{Bool}}$
$\frac{\Gamma \vdash b : \text{Bool} \quad \Gamma \vdash e_1 : t \quad \Gamma \vdash e_2 : t}{\Gamma \vdash \text{If}(b, e_1, e_2) : t}$
$\overline{\Gamma \vdash \text{Var}(x) : t} \text{ si } x : t \in \Gamma$
$\overline{\Gamma \vdash \text{Cte}(n) : \text{Int}}$
$\frac{\Gamma \vdash e_1 : \text{Int} \quad \Gamma \vdash e_2 : \text{Int}}{\Gamma \vdash \times (e_1, e_2) : \text{Int}}$
$\frac{\Gamma \vdash e_1 : \text{Int} \quad \Gamma \vdash e_2 : \text{Int}}{\Gamma \vdash + (e_1, e_2) : \text{Int}}$
$\frac{\Gamma \vdash e_1 : t_1 \quad \Gamma \cdot \{x : t_1\} \vdash e_2 : t_2}{\Gamma \vdash \text{Let}(x, e_1, e_2) : t_2}$

FIG. 3.6 – *Typage des expressions arithmético-logiques étendues.*

$$\begin{array}{c}
\overline{\Gamma \vdash \text{True} \Rightarrow T} \\
\\
\overline{\Gamma \vdash \text{False} \Rightarrow F} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{And}(e_1, e_2) \Rightarrow v_1 \wedge v_2} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{Or}(e_1, e_2) \Rightarrow v_1 \vee v_2} \\
\\
\frac{\Gamma \vdash e \Rightarrow v}{\Gamma \vdash \text{Not}(e) \Rightarrow \neg v} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash = (e_1, e_2) \Rightarrow T} \text{ si } v_1 = v_2 \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash = (e_1, e_2) \Rightarrow F} \text{ si } v_1 \neq v_2 \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash > (e_1, e_2) \Rightarrow T} \text{ si } v_1 > v_2 \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash > (e_1, e_2) \Rightarrow F} \text{ si } v_1 \not> v_2 \\
\\
\frac{\Gamma \vdash b \Rightarrow T \quad \Gamma \vdash e_1 \Rightarrow v_1}{\Gamma \vdash \text{If}(b, e_1, e_2) \Rightarrow v_1} \\
\\
\frac{\Gamma \vdash b \Rightarrow F \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{If}(b, e_1, e_2) \Rightarrow v_2} \\
\\
\overline{\Gamma \vdash \text{Var}(x) \Rightarrow v} \text{ si } x = v \in \Gamma \\
\\
\overline{\Gamma \vdash \text{Cte}(n) \Rightarrow n} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v \quad \Gamma \vdash e_2 \Rightarrow v'}{\Gamma \vdash \times (e_1, e_2) \Rightarrow v \times v'} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v \quad \Gamma \vdash e_2 \Rightarrow v'}{\Gamma \vdash + (e_1, e_2) \Rightarrow v + v'} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \cdot \{x = v_1\} \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{Let}(x, e_1, e_2) \Rightarrow v_2}
\end{array}$$

FIG. 3.7 – Évaluation des expressions arithmético-logiques étendues

3.3 Exercices

Exercice 3.1

Pour chacune des expressions suivantes, donner sa syntaxe abstraite, la typer, puis l'évaluer :

1. `let x = 2 in let n = x + 1 in let n = n + x in n * x`
2. `let x = let x = 2 in x + 3 in let x = x + 4 in x + 5`



Exercice 3.2

Montrer que le langage des expressions arithmétiques étendues bien typés (c'est-à-dire dans lesquelles tout identificateur est déclaré avant d'être utilisé) ne peut pas être engendré par une grammaire hors-contexte.

Exercice 3.3

Le langage des polynômes à coefficients entiers sur la variable x est donné par la grammaire concrète ambiguë suivante :

$$\begin{aligned}
 \text{Poly} &\rightarrow \text{Poly} + \text{Poly} \mid \text{Poly} * \text{Poly} \mid \text{Poly} \sim \text{Nat} \\
 &\mid \text{Nat} \mid - \text{Nat} \mid \mathbf{x} \mid (\text{Poly}) \\
 \text{Nat} &\rightarrow (0 - 9)^+
 \end{aligned}$$

Pour résoudre les ambiguïtés, on convient que le $+$ et le $*$ associent à gauche, et que $\sim$ est prioritaire sur $*$ lui-même prioritaire sur $+$. Ainsi, le polynôme $P_0 = (x + 1)(2x^2 - 1)$ pourra être écrit $(\mathbf{x}+1)*(\mathbf{2}*\mathbf{x}^2+\mathbf{-1})$.

Pour représenter abstraitement les polynômes, on définit la grammaire abstraite suivante :

$$\begin{aligned}
 \text{Poly} &\rightarrow +(\text{Poly} , \text{Poly}) \mid \times (\text{Poly} , \text{Poly}) \mid \text{Puiss}(\text{Poly} , \text{Int}) \\
 &\mid \text{Cte}(\text{Int}) \mid \mathbf{X} \\
 \text{Int} &\rightarrow (0 \mid \dots \mid 9)^+ \mid -(0 \mid \dots \mid 9)^+
 \end{aligned}$$

1. Énumérez les étiquettes abstraites de cette grammaire.
2. Donner l'expression abstraite représentant P_0 .

Correction : $\times(+(X, \text{Cte}(1)), +(\times(\text{Cte}(2), \text{Puiss}(X, 2)), \text{Cte}(-1)))$

3. Donner une grammaire non ambiguë pour le langage des polynômes, respectant les priorités. Associer à chaque règle une action qui construit la représentation abstraite associée.

Correction :

$$\begin{aligned}
 \text{Poly} &\rightarrow \text{Poly} + Q && \{+(\text{Poly}_1, Q_1)\} \\
 \text{Poly} &\rightarrow Q && \{ Q_1 \} \\
 Q &\rightarrow Q \times R && \{ \times(Q_1, R_1) \} \\
 Q &\rightarrow R && \{ R_1 \} \\
 R &\rightarrow R \sim \text{Nat} && \{ \text{Puiss}(R_1, \text{val}(\text{Nat}_1)) \} \\
 R &\rightarrow S && \{ S_1 \} \\
 S &\rightarrow (\text{Poly}) && \{ \text{Poly}_1 \} \\
 S &\rightarrow \text{Nat} && \{ \text{Cte}(\text{val}(\text{Nat}_1)) \} \\
 S &\rightarrow - \text{Nat} && \{ \text{Cte}(-\text{val}(\text{Nat}_1)) \} \\
 S &\rightarrow \mathbf{x} && \{ X \}
 \end{aligned}$$

4. On désire, à l'aide de règles sémantiques, effectuer le calcul de la dérivée d'un polynôme. On considère des jugements de la forme $P \Rightarrow Q$, où P et Q sont des expressions abstraites de polynômes, signifiant « Q est la dérivée de P ». Définir les règles sémantiques nécessaires.

Correction : Notons que l'exposant d'un polynôme ne peut qu'être positif ou nul après la phase d'analyse. Il n'est donc pas besoin de donner les règles pour le cas d'un exposant négatif.

$$\begin{array}{c}
 \frac{e \Rightarrow e' \quad f \Rightarrow f'}{+(e, f) \Rightarrow +(e', f')} \quad \frac{e \Rightarrow e' \quad f \Rightarrow f'}{\times(e, f) \Rightarrow (\times(e, f'), \times(e', f))} \\
 \frac{e \Rightarrow e'}{\text{Puiss}(e, n) \Rightarrow (\times(\text{Cte}(n), \times(e', \text{Puiss}(e, n-1))))} \text{ si } n > 1 \\
 \frac{e \Rightarrow e'}{\text{Puiss}(e, 1) \Rightarrow e'} \quad \frac{}{\text{Puiss}(e, 0) \Rightarrow \text{Cte}(0)} \\
 \frac{}{\text{Cte}(n) \Rightarrow \text{Cte}(0)} \quad \frac{}{X \Rightarrow \text{Cte}(1)}
 \end{array}$$

5. Donner la dérivation du jugement $a \Rightarrow P$ où a est l'arbre abstrait de $((x^3)+(2*x))$ et P est le polynôme approprié.

Correction :

$$\begin{array}{c}
 \frac{X \Rightarrow \text{Cte}(1)}{\text{Puiss}(X, 3) \Rightarrow \times(\text{Cte}(3), \times(\text{Cte}(1), \text{Puiss}(X, 2)))} \quad \frac{\frac{\text{Cte}(2) \Rightarrow \text{Cte}(0)}{\times(\text{Cte}(2), X) \Rightarrow \times(\text{Cte}(0), X)}, \frac{X \Rightarrow \text{Cte}(1)}{\times(\text{Cte}(2), \text{Cte}(1))}}{+(\times(\text{Cte}(0), X), \times(\text{Cte}(2), \text{Cte}(1)))} \\
 \frac{+(\text{Puiss}(X, 3), \times(\text{Cte}(2), X)) \Rightarrow}{+(\times(\text{Cte}(3), \text{Puiss}(X, 2)), +(\times(\text{Cte}(0), X), \times(\text{Cte}(2), \text{Cte}(1))))}
 \end{array}$$

On désire rajouter dans le langage la possibilité de définir des variables associées à des polynômes à l'aide de la construction `let ... in ...`, comme par exemple

`let P = x^2+3*x+1 in let Q = P^2+x in (2*x*P+Q)^3`

Les noms de variables seront formés d'une lettre majuscule suivie éventuellement de chiffres.

1. Donner les règles supplémentaires à ajouter dans la syntaxe concrète et la syntaxe abstraite.

Correction : Nouvelles règles de syntaxe concrète :

$$\begin{array}{lcl}
 P & \rightarrow & \text{let } Id = P \text{ in } P \quad \{ \text{Let } (\text{val}(Id_1), P_1, P_2) \} \\
 P & \rightarrow & Poly \quad \{ Poly_1 \} \\
 S & \rightarrow & Id \quad \{ \text{Var } (\text{val}(Id_1)) \}
 \end{array}$$

Nouvelles règles de syntaxe abstraite :

$$Poly \rightarrow \text{Let } (String, Poly, Poly) \mid \text{Var } (String)$$

2. Définir un nouveau système de règles sémantiques pour le calcul de la dérivée, s'appliquant sur un jugement de la forme appropriée. Attention, pour le calcul de la dérivée d'une expression polynomiale contenant un `let`, on demande de préserver le `let` : par exemple, la dérivée de `let P=x in P^2` sera `let P=x in 2*P^1*1`.

Correction : Attention : il faut calculer la dérivée du polynôme e figurant dans $\text{Let}(P, e, f)$ dans le bon environnement. Cela impose de la calculer immédiatement, et les environnements

seront donc de la forme $\{P_1 = e_1, \dots\}$, où e_1 désigne la dérivée du polynôme valeur de P , et non pas le polynôme valeur de P lui-même.

$$\begin{array}{c}
\frac{\Gamma \vdash e \Rightarrow e' \quad \Gamma \cdot \{P = e'\} \vdash f \Rightarrow f'}{\Gamma \vdash \text{Let } (P, e, f) \Rightarrow \text{Let } (P, e, f')} \quad \frac{}{\Gamma \vdash \text{Var } (P) \Rightarrow e} \text{ si } P = e \in \Gamma \\
\\
\frac{\Gamma \vdash e \Rightarrow e' \quad \Gamma \vdash f \Rightarrow f'}{\Gamma \vdash + (e, f) \Rightarrow + (e', f')} \quad \frac{\Gamma \vdash e \Rightarrow e' \quad \Gamma \vdash f \Rightarrow f'}{\Gamma \vdash \times (e, f) \Rightarrow + (\times (e, f'), \times (e', f))} \\
\\
\frac{\Gamma \vdash e \Rightarrow e'}{\Gamma \vdash \text{Puiss } (e, n) \Rightarrow \times (\text{Cte } (n), \times (e', \text{Puiss } (e, n - 1)))} \text{ si } n > 1 \\
\\
\frac{\Gamma \vdash e \Rightarrow e'}{\Gamma \vdash \text{Puiss } (e, 1) \Rightarrow e'} \quad \frac{}{\Gamma \vdash \text{Puiss}(e, 0) \Rightarrow \text{Cte } (0)} \\
\\
\frac{}{\Gamma \vdash \text{Cte } (n) \Rightarrow \text{Cte } (0)} \quad \frac{}{\Gamma \vdash X \Rightarrow \text{Cte } (1)}
\end{array}$$

3. On suppose maintenant que les polynômes sont écrits sous la forme d'une somme de monômes, et l'on se pose les même questions que précédemment, à savoir, la définition d'une syntaxe concrète non-ambiguë, d'une syntaxe abstraite, de règles de mise sous forme d'une somme de monômes pour une somme et un produit de polynômes, et de règles de calcul de la dérivée d'un polynôme.

Chapitre 4

Fonctions à un argument

Ce chapitre s'intéresse aux définitions de fonctions dans un langage de programmation.

Pour conserver un langage simple, nous allons supposer que nous définissons des fonctions locales à une expression arithmétique grâce à une construction `let`. Définir des fonctions « globales » ne poserait pas de difficulté supplémentaire, car on peut toujours considérer qu'une fonction globale est une fonction locale à tout le programme qui la suit.

Pour éviter d'alourdir la syntaxe, nous considérerons dans ce chapitre uniquement des fonctions à un argument. Définir des fonctions à plusieurs arguments peut se faire par les mêmes méthodes, au prix d'une complication de syntaxe, ceci sera l'objet de l'exercice 4.2.

Nous commençons par le cas des fonctions non récursives, puis le cas de la récursivité sera traité spécialement.

4.1 Fonctions non récursives

La syntaxe concrète que nous adoptons pour définir des fonctions est de la forme `let fun  $f(x : t_1) : t_2 = e_1$  in  $e_2$` signifiant que l'on définit une fonction de *nom* f , de *paramètre formel* x de type t_1 , dont le type du résultat est t_2 . e_1 et e_2 sont des expressions quelconques, e_1 est le *corps* de la fonction et e_2 est l'expression dans laquelle f est utilisée. Nous supposons que les types sont uniquement des types simples classiques : entiers, booléens, réels, chaînes de caractères, etc. Dans ce chapitre nous prenons les entiers et les booléens pour fixer les idées.

Enfin, l'application d'une fonction f à une expression e sera notée classiquement par $f(e)$.

4.1.1 Syntaxe abstraite

La syntaxe abstraite choisie pour les fonctions à un argument est donnée par la grammaire des expressions arithmético-logiques avec variables à laquelle on rajoute

$$\begin{aligned} E &\rightarrow \text{LetFun}(\text{String}, \text{String}, T, T, E, E) \mid \text{App}(\text{String}, E) \\ T &\rightarrow \text{Int} \mid \text{Bool} \mid \dots \end{aligned}$$

$$\boxed{
\begin{array}{c}
\frac{\Gamma \cdot \{x : t_1\} \vdash e_1 : t_2 \quad \Gamma \cdot \{f : t_1 \rightarrow t_2\} \vdash e_2 : t_3}{\Gamma \vdash \text{LetFun}(f, x, t_1, t_2, e_1, e_2) : t_3} \\
\\
\frac{\Gamma \vdash e : t_1}{\Gamma \vdash \text{App}(f, e) : t_2} \text{ si } f : t_1 \rightarrow t_2 \in \Gamma
\end{array}
}$$

FIG. 4.1 – *Typage des expressions arithmétiques fonctionnelles*

La grammaire concrète non-ambiguë accompagnée des règles de calcul de la syntaxe abstraite est alors :

$$\begin{array}{ll}
E \rightarrow \text{let } Id = E \text{ in } E & \{\text{Let}(\text{val}(Id_1), E_1, E_2)\} \\
E \rightarrow \text{let fun } Id(Id : T) : T = E \text{ in } E & \{\text{LetFun}(\text{val}(Id_1), \text{val}(Id_2), \text{val}(T_1), \text{val}(T_2), E_1, E_2)\} \\
E \rightarrow F & \{F_1\} \\
F \rightarrow F + G & \{+(F_1, G_1)\} \\
F \rightarrow G & \{G_1\} \\
G \rightarrow G * H & \{\times(G_1, H_1)\} \\
G \rightarrow H & \{H_1\} \\
H \rightarrow \text{Nat} & \{\text{Cte}(\text{val}(\text{Nat}_1))\} \\
H \rightarrow Id & \{\text{Var}(\text{val}(Id_1))\} \\
H \rightarrow Id(E) & \{\text{App}(\text{val}(Id_1), E_1)\} \\
H \rightarrow (E) & \{E_1\} \\
T \rightarrow \text{int} & \{\text{Int}\} \\
T \rightarrow \text{bool} & \{\text{Bool}\}
\end{array}$$

Par exemple, à l'expression concrète `let fun sqr(x : int) : int = x × x in sqr(3) + sqr(4)` correspond la syntaxe abstraite

$$\text{LetFun}(\text{sqr}, x, \text{Int}, \text{Int}, \times(\text{Var}(x), \text{Var}(x)), +(\text{App}(\text{sqr}, \text{Cte}(3)), \text{App}(\text{sqr}, \text{Cte}(4))))$$

4.1.2 Typage

Nos jugements de typage s'écriront encore $\Gamma \vdash e : t$, où e dénote une expression de syntaxe abstraite. Mais ici, l'ensemble des types est plus riche : nous aurons les types des fonctions des entiers ou les booléens vers les entiers ou les booléens. En effet, une fonction des entiers naturels dans les entiers naturels est un objet très différent des identificateurs de type entier. En particulier, un identificateur de fonction doit toujours être utilisé avec un argument, de type entier dans le cas qui nous intéresse. On dit qu'un identificateur de fonction a un *type fonctionnel*, que l'on note $t_1 \rightarrow t_2$ si l'argument est de type t_1 et le résultat est de type t_2 .

Les règles de typage sont décrites sur la figure 4.1. Les anciennes règles de la figure 3.6 sont encore valables et ne sont pas répétées.

Notons à nouveau la règle de portée : dans la définition `let fun f(x : t1) : t2 = e1 in e2`, le nom x du paramètre formel peut être utilisé dans e_1 mais pas dans e_2 , alors que le nom f de la fonction peut être utilisé dans e_2 mais pas dans e_1 , il ne peut donc pas s'agir d'une fonction récursive, cas que nous allons traiter un peu plus loin.

La dérivation de typage de l'exemple de `sqr` ci-dessus est donné sur la figure 4.2.

4.1.3 Évaluation

Ces expressions arithmétiques fonctionnelles s'évaluent de manière similaire aux expressions arithmétiques usuelles, mais il va falloir cette fois un environnement liant le nom

$$\begin{array}{c}
(1) : \frac{\frac{\Gamma \vdash \text{Cte}(3) : \text{Int}}{\Gamma \vdash \text{App}(sqr, \text{Cte}(3)) : \text{Int}} \quad \frac{\Gamma \vdash \text{Cte}(4) : \text{Int}}{\Gamma \vdash \text{App}(sqr, \text{Cte}(4)) : \text{Int}}}{\Gamma \vdash \times(\text{App}(sqr, \text{Cte}(3)), \text{App}(sqr, \text{Cte}(4))) : \text{Int}} \\
\\
\frac{\frac{\{x : \text{Int}\} \vdash \text{Var}(x) : \text{Int}}{\{x : \text{Int}\} \vdash \times(\text{Var}(x), \text{Var}(x)) : \text{Int}} \quad (1)}{\{\} \vdash \text{LetFun}(sqr, x, \times(\text{Var}(x), \text{Var}(x)), +(\text{App}(sqr, \text{Cte}(3)), \text{App}(sqr, \text{Cte}(4)))) : \text{Int}}
\end{array}$$

FIG. 4.2 – Typage de l'expression arithmétique fonctionnelle *let fun* $sqr(x) = x \times x$ *in* $sqr(3) + sqr(4)$ où $\Gamma = \{sqr : \text{Int} \rightarrow \text{Int}\}$

$$\begin{array}{c}
\frac{\Gamma \cdot \{f = (x, e_1)\} \vdash e_2 \Rightarrow v}{\Gamma \vdash \text{LetFun}(f, x, e_1, e_2) \Rightarrow v} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \cdot \{x = v_1\} \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{App}(f, e_1) \Rightarrow v_2} \text{ si } f = (x, e_2) \in \Gamma
\end{array}$$

FIG. 4.3 – Évaluation des expressions arithmétiques fonctionnelles

d'une fonction à son corps, c'est-à-dire à une expression arithmétique fonctionnelle. De plus, le corps dépend du paramètre formel de la fonction, paramètre qui sera lié lors de l'appel à la fonction. Il va donc falloir stocker dans l'environnement de liaison aussi bien le corps que le nom du paramètre formel de la fonction. On utilisera pour cela la notation $f = (x, e)$, où f est le nom de la fonction, x est le nom de son paramètre formel, et e est son corps. Sur l'exemple précédent l'environnement d'évaluation où sqr est définie est noté $\{sqr = (x, \times(\text{Var}(x), \text{Var}(x)))\}$.

Le domaine de valeur des jugements d'évaluation sera donc $D^v = \mathbb{N} \cup \{T, F\}$ et le domaine de contexte D^c sera l'ensemble des listes de couples (identificateur, valeur) où les valeurs sont soit des entiers ou des booléens pour les identificateurs de variables, soit des couples (identificateur, expression abstraite) pour les identificateurs de fonctions.

Les règles d'évaluation sont données sur la figure 4.3, les anciennes règles de la figure 3.7 étant encore valables.

Théorème 4.1 *Le langage des expressions arithmétiques fonctionnelles, avec la sémantique définie ci-dessus, n'est pas fortement typé. Voir exercice 4.4.*

La dérivation d'évaluation de notre exemple est traité sur la figure 4.4, où sur cette figure Γ dénote l'environnement $\{sqr = (x, \times(\text{Var}(x), \text{Var}(x)))\}$. D'autres exemples instructifs de typage et d'évaluation sont proposés en exercice. Notons que les calculs d'évaluation terminent toujours pour tous les langages définis jusqu'à présent.

4.2 Fonctions récursives

Nous avons traité dans le paragraphe précédent des fonctions non récursives, nous allons maintenant considérer le cas plus complexe des fonctions récursives.

Dans la syntaxe concrète nous choisissons de rajouter le mot clé **rec** pour spécifier que l'on définit une fonction récursive, et dans la syntaxe abstraite nous introduisons une construction *LetFunRec* analogue au *LetFun*.

$$\begin{array}{c}
(1) : \frac{\frac{\text{trivial} \dots}{\Gamma \vdash \text{Cte}(3) \Rightarrow 3} \quad \frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3\} \vdash \times(\text{Var}(x), \text{Var}(x)) \Rightarrow 9}}{\Gamma \vdash \text{App}(sqr, \text{Cte}(3)) \Rightarrow 9} \\
(2) : \frac{\frac{\text{trivial} \dots}{\Gamma \vdash \text{Cte}(4) \Rightarrow 4} \quad \frac{\text{trivial} \dots}{\Gamma \cdot \{x = 4\} \vdash \times(\text{Var}(x), \text{Var}(x)) \Rightarrow 16}}{\Gamma \vdash \text{App}(sqr, \text{Cte}(4)) \Rightarrow 16} \\
\frac{(1) \quad (2)}{\Gamma \vdash \text{App}(sqr, \text{Cte}(3)) \Rightarrow 9 \quad \Gamma \vdash \text{App}(sqr, \text{Cte}(4)) \Rightarrow 16} \\
\frac{\Gamma \vdash \text{App}(sqr, \text{Cte}(3)) \Rightarrow 9 \quad \Gamma \vdash \text{App}(sqr, \text{Cte}(4)) \Rightarrow 16}{\Gamma \vdash +(\text{App}(sqr, \text{Cte}(3)), \text{App}(sqr, \text{Cte}(4))) \Rightarrow 25} \\
\frac{\Gamma \vdash +(\text{App}(sqr, \text{Cte}(3)), \text{App}(sqr, \text{Cte}(4))) \Rightarrow 25}{\{\} \vdash \text{LetFun}(sqr, n, \times(\text{Var}(x), \text{Var}(x)), +(\text{App}(sqr, \text{Cte}(3)), \text{App}(sqr, \text{Cte}(4)))) \Rightarrow 25}
\end{array}$$

FIG. 4.4 – Évaluation de l'expression arithmétique fonctionnelle *let fun* $sqr(x) = x \times x$ *in* $sqr(3) + sqr(4)$ où $\Gamma = \{sqr = (x, \times(\text{Var}(x), \text{Var}(x)))\}$.

$$\frac{\Gamma \cdot \{f : t_1 \rightarrow t_2, x : t_1\} \vdash e_1 : t_2 \quad \Gamma \cdot \{f : t_1 \rightarrow t_2\} \vdash e_2 : t}{\Gamma \vdash \text{LetFunRec}(f, n, t_1, t_2, e_1, e_2) : t}$$

FIG. 4.5 – Typage des fonctions récursives

Nous allons considérer l'exemple fondamental de la fonction factorielle définie dans notre syntaxe par

```
let fun rec f(x : int) : int = if x = 0 then 1 else x × f(x - 1) in f(3)
```

4.2.1 Typage

Les règles de typage sont très similaires à celles des paragraphes précédents. La différence essentielle concerne la portée du nom de la fonction récursive définie par **let fun rec**. Le nom de la fonction peut en effet intervenir dans le corps de la fonction, et il faut donc ajouter le type de la fonction à l'environnement de typage avant d'en typer le corps. Cela suppose bien sûr que ce type est connu, ce qui est le cas dans le cadre que nous avons choisi¹. Nous nous contentons de donner à la figure 4.5 la nouvelle règle de typage des expressions **let fun rec**. L'exemple donné au paragraphe précédent montre un cas où le résultat est du même type que la fonction. L'exemple ci-dessous donne un exemple de fonction qui a le type `int -> int` et pourtant le type de l'expression **let fun rec** est `bool` :

```
let fun rec f(x : int) : int = if x = 0 then 1 else x × f(x - 1) in f(10) > 1000
```

Le typage de cette expression est laissé au lecteur !

4.2.2 Évaluation

Les règles d'évaluation sont inchangées pour les expressions **LetFunRec** par rapport aux expressions **LetFun**. En effet, d'une part l'information supplémentaire concernant le type de l'argument d'une fonction récursive n'est d'aucune utilité à l'évaluation ; d'autre part, la notion d'environnement que nous avons utilisée est suffisamment générale pour couvrir

¹. L'inférence de types comme dans CAML est hors du cadre de ce cours.

le cas des fonctions récursives. On remarquera néanmoins l'importance de la gestion des environnements en pile, qui permet de récupérer à chaque fois la bonne valeur du paramètre formel de la fonction récursive. Ceci est mis en évidence sur l'évaluation de l'expression calculant la factorielle de 3, qui est donnée sur les figures 4.6 et 4.7.

4.3 Exercices

Exercice 4.1

Analyser, typer, puis évaluer si possible les expressions suivantes.

1. `let fun sqr(x : int) : int = x * x in sqr(1) + sqr(sqr(2)).`
2. `let fun sqr(x : int) : int = x * x in let x = sqr(1) + sqr(sqr(2)) in sqr(x).`
3. `let fun f(x : int) : int = x * f(x) in f(1).`
4. `let fun f(x : int) : int = x * x in f(x + 1).`

Exercice 4.2

On veut généraliser l'étude de ce chapitre aux fonctions à plusieurs arguments. Pour ce faire, on modifie la règle de la syntaxe abstraite pour le LetFun de la manière suivante :

$$\begin{aligned} E &\rightarrow \text{LetFun}(\text{String}, \text{Args}, T, E, E) \\ \text{Args} &\rightarrow \text{Argvide} \mid \text{Argcons}(\text{String}, T, \text{Args}) \end{aligned}$$

1. Donner la nouvelle grammaire concrète non-ambiguë accompagnée des calculs de syntaxe abstraite.
2. Donner les nouvelles règles de typage.
3. Donner les nouvelles règles d'évaluation.
4. Donner les dérivations de typage et d'évaluation du programme `let fun f(x : int, y : int) : int = x * x + y * y in f(3, 4)`

Exercice 4.3

On veut généraliser l'étude de ce chapitre aux fonctions mutuellement récursives. Il faut pour cela être capable de définir simultanément deux fonctions, et les constructions `let fun` doivent donc être capable de traiter une liste quelconque de définitions. Nous choisissons alors la syntaxe abstraite suivante :

$$\begin{aligned} E &\rightarrow \text{LetFunRec}(\text{Defs}, E) \\ \text{Defs} &\rightarrow \text{Defvide} \mid \text{Defcons}(\text{String}, \text{String}, T, T, E, \text{Defs}) \end{aligned}$$

1. Donner la nouvelle grammaire concrète non-ambiguë accompagnée des calculs de syntaxe abstraite.
2. Donner les nouvelles règles de typage.
3. Donner les nouvelles règles d'évaluation.
4. Donner les dérivations de typage et d'évaluation du programme

```
let fun recf(x : int) : bool = if x = 0 then true else g(x - 1)
    and also g(x : int) : bool = if x = 0 then false else f(x - 1)
    in f(3)
```

$$\begin{array}{l}
(5) : \frac{\frac{\Gamma \cdot \{x = 3, x = 2\} \vdash \text{Var}(x) \Rightarrow 2}{\Gamma \cdot \{x = 3, x = 2\} \vdash \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 2} \quad \frac{\Gamma \cdot \{x = 3, x = 2\} \vdash \text{App}(f, -(\text{Var}(x), \text{Cte}(1))) \Rightarrow 1}{\Gamma \cdot \{x = 3, x = 2\} \vdash \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 2} \quad (6)}{\Gamma \cdot \{x = 3, x = 2\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))))) \Rightarrow 2} \quad (5) \\
(4) : \frac{\frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3, x = 2\} \vdash = (\text{Var}(x), \text{Cte}(0)) \Rightarrow F} \quad \frac{\Gamma \cdot \{x = 3, x = 2\} \vdash \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 2}{\Gamma \cdot \{x = 3, x = 2\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))))) \Rightarrow 2} \quad (5)}{\Gamma \cdot \{x = 3, x = 2\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))))) \Rightarrow 2} \quad (4) \\
(3) : \frac{\frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3\} \vdash -(\text{Var}(x), \text{Cte}(1)) \Rightarrow 2} \quad \frac{\Gamma \cdot \{x = 3, x = 2\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))))) \Rightarrow 2}{\Gamma \cdot \{x = 3\} \vdash \text{App}(f, -(\text{Var}(x), \text{Cte}(1))) \Rightarrow 2} \quad (4)}{\Gamma \cdot \{x = 3\} \vdash \text{App}(f, -(\text{Var}(x), \text{Cte}(1))) \Rightarrow 2} \quad (3) \\
(2) : \frac{\frac{\Gamma \cdot \{x = 3\} \vdash \text{Var}(x) \Rightarrow 3}{\Gamma \cdot \{x = 3\} \vdash \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 6} \quad \frac{\Gamma \cdot \{x = 3\} \vdash \text{App}(f, -(\text{Var}(x), \text{Cte}(1))) \Rightarrow 2}{\Gamma \cdot \{x = 3\} \vdash \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 6} \quad (3)}{\Gamma \cdot \{x = 3\} \vdash \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 6} \quad (2) \\
(1) : \frac{\frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3\} \vdash = (\text{Var}(x), \text{Cte}(0)) \Rightarrow F} \quad \frac{\Gamma \cdot \{x = 3\} \vdash \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 6}{\Gamma \cdot \{x = 3\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))))) \Rightarrow 6} \quad (2)}{\Gamma \cdot \{x = 3\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))))) \Rightarrow 6} \quad (1) \\
\frac{\frac{\frac{\{\} \vdash \text{Cte}(3) \Rightarrow 3}{\Gamma = \{f = (x, \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1))))))\} \vdash \text{App}(f, \text{Cte}(3)) \Rightarrow 6} \quad \frac{\Gamma \cdot \{x = 3\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))))) \Rightarrow 6}{\{\} \vdash \text{LetFunRec}(f, n, \text{Int}, \text{Int}, \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1))))), \text{App}(f, \text{Cte}(3)))) \Rightarrow 6} \quad (1)}{\{\} \vdash \text{LetFunRec}(f, n, \text{Int}, \text{Int}, \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times(\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1))))), \text{App}(f, \text{Cte}(3)))) \Rightarrow 6} \quad (1)
\end{array}$$

FIG. 4.6 – *Évaluation de l'expression fonctionnelle récursive let fun rec f(x : int) : int = if x = 0 then 1 else x × f(x - 1) in f(3), partie 1*

$$\begin{array}{lcl}
(10) : & \frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3, x = 2, x = 1, x = 0\} \vdash = (\text{Var}(x), \text{Cte}(0)) \Rightarrow T} & \frac{\Gamma \cdot \{x = 3, x = 2, x = 1, x = 0\} \vdash \text{Cte}(1) \Rightarrow 1}{\Gamma \cdot \{x = 3, x = 2, x = 1, x = 0\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times (\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1))))) \Rightarrow 1} \\
(9) : & \frac{\frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash -(\text{Var}(x), \text{Cte}(1)) \Rightarrow 0} \quad \frac{(10)}{\Gamma \cdot \{x = 3, x = 2, x = 1, x = 0\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times (\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1))))) \Rightarrow 1}}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash \text{App}(f, -(\text{Var}(x), \text{Cte}(1))) \Rightarrow 1} \\
(8) : & \frac{\frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash \text{Var}(x) \Rightarrow 1} \quad \frac{(9)}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash \text{App}(f, -(\text{Var}(x), \text{Cte}(1))) \Rightarrow 1}}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash \times (\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 1} \\
(7) : & \frac{\frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash = (\text{Var}(x), \text{Cte}(0)) \Rightarrow F} \quad \frac{(8)}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash \times (\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1)))) \Rightarrow 1}}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times (\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1))))) \Rightarrow 1} \\
(6) : & \frac{\frac{\text{trivial} \dots}{\Gamma \cdot \{x = 3, x = 2\} \vdash -(\text{Var}(x), \text{Cte}(1)) \Rightarrow 1} \quad \frac{(7)}{\Gamma \cdot \{x = 3, x = 2, x = 1\} \vdash \text{If}(= (\text{Var}(x), \text{Cte}(0)), \text{Cte}(1), \times (\text{Var}(x), \text{App}(f, -(\text{Var}(x), \text{Cte}(1))))) \Rightarrow 1}}{\Gamma \cdot \{x = 3, x = 2\} \vdash \text{App}(f, -(\text{Var}(x), \text{Cte}(1))) \Rightarrow 1}
\end{array}$$

FIG. 4.7 – *Évaluation de l'expression fonctionnelle récursive let fun rec f(x : int) : int = if x = 0 then 1 else x × f(x - 1) in f(3), partie 2*

Exercice 4.4

La sémantique définie par les règles données dans ce chapitre est la sémantique par *liaison dynamique*. Les règles de sémantique pour la *liaison statique* sont :

$$\frac{\Gamma \cdot \{f = (x, e_1, \Gamma)\} \vdash e_2 \Rightarrow v}{\Gamma \vdash \text{LetFun}(f, x, e_1, e_2) \Rightarrow v}$$

$$\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma' \cdot \{x = v_1\} \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{App}(f, e_1) \Rightarrow v_2} \text{ si } f = (x, e_2, \Gamma') \in \Gamma$$

1. Évaluer le programme suivant pour chacune des deux sémantiques :

```
let x=1 in
let fun f(y:int):int = x+y in
let x=2 in f(3)
```

2. Commenter les résultats obtenus.

Chapitre 5

Sémantique d'un langage fonctionnel

Notre traitement des expressions arithmétiques permettait des définitions de fonctions à un argument. Nous pourrions étendre ce traitement de manière à traiter des fonctions à n arguments. Il nous faudrait pour cela ajouter deux nouvelles règles de typage, et deux nouvelles règles sémantiques pour chaque valeur de n . Nous ne pourrions donc traiter qu'un nombre fini de valeurs particulières de n , sauf à utiliser des règles paramétrées par n . De plus, nos fonctions seraient à valeur dans un type de base, entier ou réel.

De fait, ce qui caractérise les langages dits *fonctionnels* est que les fonctions sont des valeurs au même titre que les entiers, les réels ou les booléens. En particulier, les arguments d'une fonction peuvent être des fonctions, et le résultat peut lui aussi être une fonction.

En fait, ce qui distinguera une fonction d'une valeur de base comme un entier ou un booléen sera son type : on dira qu'une fonction est une expression de type *fonctionnel*, un type fonctionnel étant un type de la forme $t_1 \rightarrow t_2$ désignant le type des fonctions de t_1 vers t_2 . L'ensemble des types de notre langage sera décrit par la grammaire de types suivante :

$$T \rightarrow \text{int} \mid \text{bool} \mid T \rightarrow T$$

Cette grammaire est ambiguë car un type de la forme $t_1 \rightarrow t_2 \rightarrow t_3$ peut être reconnu sous les deux formes différentes $(t_1 \rightarrow t_2) \rightarrow t_3$ et $t_1 \rightarrow (t_2 \rightarrow t_3)$, or ces deux types sont tout à fait différents puisque le premier dénote les fonctions dont l'argument est une fonction de t_1 dans t_2 , alors que le second dénote les fonctions dont le résultat est une fonction de t_2 dans t_3 . Traditionnellement, on considère que la flèche associée à droite, c'est-à-dire que la seconde forme $t_1 \rightarrow (t_2 \rightarrow t_3)$ est la forme par défaut. Pour spécifier une fonction dont le type est $(t_1 \rightarrow t_2) \rightarrow t_3$, on utilisera des parenthèses. La grammaire non-ambiguë des types de notre langage est donc

$$\begin{aligned} T &\rightarrow T' \rightarrow T \mid T' \\ T' &\rightarrow \text{int} \mid \text{bool} \mid (T) \end{aligned}$$

Dans ce chapitre, nous allons dans un premier temps décrire deux primitives à la base de tous les calculs fonctionnels, l'abstraction et l'application. Puis nous verrons comment coder les fonctions à n arguments grâce à ces nouvelles primitives, ainsi que les instructions *let* déjà vues.

5.1 Syntaxes concrète et abstraite

La fonction qui à la variable x associe le calcul décrit par l'expression e est appelée une *abstraction* et est notée (il s'agit de la syntaxe concrète) $\text{fun } x : t \rightarrow e$. Cela correspond à la notation mathématique $x \mapsto f(x)$, et permet de définir une fonction qui n'a pas de nom. L'expression e est appelée le *corps* de la fonction.

L'appel fonctionnel sera maintenant remplacé par l'application d'une abstraction $\text{fun } x : t \rightarrow e_1$ à un argument e_2 . Cette application sera notée (en syntaxe concrète) $(\text{fun } x : t \rightarrow e_1) e_2$. Plus généralement, si e_1 et e_2 sont deux expressions, on notera l'application de e_1 à e_2 par $(e_1 e_2)$, les parenthèses pouvant être omises s'il n'y a aucune ambiguïté. La syntaxe concrète de nos expressions est donc décrite par la grammaire YACC non-ambiguë ci-dessous, qui calcule en même temps la syntaxe abstraite des expressions :

$E \rightarrow \text{fun } Id : T \rightarrow E$	$\{\text{Fun}(\text{val}(Id_1), T_1, E_1)\}$
$E \rightarrow \text{let } Id = E \text{ in } E$	$\{\text{Let}(\text{val}(Id_1), E_1, E_1)\}$
$E \rightarrow \text{if } E \text{ then } E \text{ else } E$	$\{\text{If}(E_1, E_2, E_1)\}$
$E \rightarrow F$	$\{F_1\}$
$F \rightarrow F \text{ or } G$	$\{\text{Or}(F_1, G_1)\}$
$F \rightarrow G$	$\{G_1\}$
$G \rightarrow G \text{ and } H$	$\{\text{And}(G_1, H_1)\}$
$G \rightarrow H$	$\{H_1\}$
$H \rightarrow \text{not } H$	$\{\text{Not}(H_1)\}$
$H \rightarrow I$	$\{I_1\}$
$I \rightarrow I = J$	$\{=(I_1, J_1)\}$
$I \rightarrow I > J$	$\{>(I_1, J_1)\}$
$I \rightarrow J$	$\{J_1\}$
$J \rightarrow J + K$	$\{+(J_1, K_1)\}$
$J \rightarrow K$	$\{K_1\}$
$K \rightarrow K * L$	$\{\times(K_1, L_1)\}$
$K \rightarrow L$	$\{L_1\}$
$L \rightarrow \text{true}$	$\{\text{True}\}$
$L \rightarrow \text{false}$	$\{\text{False}\}$
$L \rightarrow \text{Nat}$	$\{\text{Cte}(\text{val}(\text{Nat}_1))\}$
$L \rightarrow Id$	$\{\text{Var}(\text{val}(Id_1))\}$
$L \rightarrow (E)$	$\{E_1\}$
$L \rightarrow (E E)$	$\{\text{App}(E_1, E_2)\}$
$T \rightarrow T' \rightarrow T$	$\{\text{Fleche}(T_1, T_2)\}$
$T' \rightarrow (T)$	$\{T_1\}$
$T' \rightarrow \text{int}$	$\{\text{Int}\}$
$T' \rightarrow \text{bool}$	$\{\text{Bool}\}$

5.2 Typage

Les expressions de syntaxe abstraites contiennent des informations qui n'ont pour but que de faciliter la phase de contrôle de type. Comme elles ne sont d'aucune utilité pour la phase d'évaluation, on va pouvoir nettoyer les expressions de syntaxe abstraite lors de la phase de typage, de façon à obtenir de nouvelles expressions de syntaxe abstraite plus simples. Pour cela, nous utiliserons une syntaxe des jugements analogue à celle du paragraphe 2.4.2. Les règles de typage de notre langage fonctionnel sont données à la figure 5.1.

$$\boxed{
\begin{array}{c}
\frac{\Gamma \cdot \{x : t_1\} \vdash e_1 : t_2 \Rightarrow e_2}{\Gamma \vdash \text{Fun}(x, t_1, e_1) : \text{Fleche}(t_1, t_2) \Rightarrow \text{Fun}(x, e_2)} \\
\\
\frac{\Gamma \vdash e_1 : \text{Fleche}(t_1, t_2) \Rightarrow e'_1 \quad \Gamma \vdash e_2 : t_1 \Rightarrow e'_2}{\Gamma \vdash \text{App}(e_1, e_2) : t_2 \Rightarrow \text{App}(e'_1, e'_2)}
\end{array}
}$$

FIG. 5.1 – *Typage des expressions fonctionnelles*

$$\boxed{
\begin{array}{c}
\overline{\Gamma \vdash \text{Fun}(x, e) \Rightarrow (x, e)} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow (x, e) \quad \Gamma \vdash e_2 \Rightarrow v_2 \quad \Gamma \cdot \{x = v_2\} \vdash e \Rightarrow v_3}{\Gamma \vdash \text{App}(e_1, e_2) \Rightarrow v_3}
\end{array}
}$$

FIG. 5.2 – *Évaluation des expressions fonctionnelles*

5.3 Évaluation

Les expressions fonctionnelles sont maintenant des expressions comme les autres, elles doivent donc avoir une valeur. En première approximation, la valeur sera un couple formé par le nom de la variable et l'expression définissant cette fonction. Il s'agit tout simplement du type d'information stocké dans l'environnement dans le cas de l'évaluation des fonctions à un argument, voir figure 4.3.

Les règles d'évaluation sont alors celles de la figure 5.2.

Par exemple, l'expression suivante

```
let succ = fun x:int -> x+1 in (succ 2)
```

conduit à évaluer d'abord l'expression $\text{App}(\text{Var}(\text{succ}), \text{Cte}(2))$ dans l'environnement $\Gamma_1 = \{\text{succ} = (x, +(\text{Var}(x), \text{Cte}(1)))\}$, puis à évaluer dans l'environnement $\Gamma_1 = \Gamma_1 \cdot \{x = 2\}$ l'expression $+(\text{Var}(x), \text{Cte}(1))$ qui donne bien le résultat escompté.

Un exemple de fonction avec fonction en argument est donné par

```
let eval0 = fun f:(int->int) -> (f 0)
in (eval0 (fun x:int -> x+1))
```

qui conduira à évaluer dans l'environnement $\Gamma = \{f = (x, +(\text{Var}(x), \text{Cte}(1)))\}$ l'expression $\text{App}(\text{Var}(f), \text{Cte}(0))$ qui retournera à son tour la valeur escomptée 1.

5.4 Fonctions à plusieurs arguments

Une fonction étant une expression parmi d'autres, une fonction à deux arguments est une abstraction $\text{fun } x \rightarrow e_1$ dont le corps est une fonction à un argument $e_1 = \text{fun } y \rightarrow e_2$. La fonction s'écrit donc $\text{fun } x \rightarrow \text{fun } y \rightarrow e_2$. On peut bien sûr tout aussi bien définir des fonctions d'un nombre quelconque d'arguments. Afin de faciliter la lecture des fonctions de multiples arguments, on a l'habitude de grouper les variables sous les `fun` successifs, avec l'écriture $\text{fun } x_1 \cdots x_n \rightarrow e$, qui est par définition une abbréviation de l'écriture $\text{fun } x_1 \rightarrow \text{fun } x_2 \rightarrow \cdots \text{fun } x_n \rightarrow e$. Cette convention d'écriture cache donc le fait qu'une fonction de deux arguments x et y et de corps e n'est rien d'autre qu'une fonction de l'unique argument x dont le corps est la fonction $\text{fun } y \rightarrow e$. Grâce à notre nouvelle notation, les fonctions de n arguments sont donc devenues des fonctions de un argument.

$$\boxed{
\begin{array}{c}
\frac{\Gamma \vdash e \Rightarrow e' \quad \Gamma \vdash g \Rightarrow g'}{\Gamma \vdash \text{Let}(n, e, g) \Rightarrow \text{App}(\text{Fun}(n, g'), e')} \\
\\
\frac{\Gamma \vdash e \Rightarrow e' \quad \Gamma \vdash g \Rightarrow g'}{\Gamma \vdash \text{LetFun}(\text{Var}(f), \text{Var}(n), e, g) \Rightarrow \text{App}(\text{Fun}(f, g'), \text{Fun}(n, e'))}
\end{array}
}$$

FIG. 5.3 – Codages des expressions *let*

5.5 Problème de la liaison dynamique

Le problème de la liaison dynamique entrevu dans l'exercice 4.4 se présente ici de manière cruciale. Par exemple considérons l'expression

```
let f = fun x:int -> fun y:int -> 2*x+3*y
in ((f 2) 3)
```

Avec les règles ci-dessus, l'évaluation de `(f 2)` va demander d'évaluer l'expression `fun y:int -> 2*x+3*y` dans l'environnement $\{x = 2\}$ qui donne le résultat $(y, 2*x+3*y)$. Mais alors l'évaluation de `((f 2) 3)` va demander d'évaluer l'expression `2*x+3*y` dans l'environnement $\{y = 3\}$ qui va échouer car x n'y a pas de valeur. Pour évaluer correctement l'expression, il est nécessaire de procéder comme dans l'exercice 4.4, en liant la fonction à l'environnement dans lequel elle est définie. La valeur d'une expression de type fonctionnel sera donc un triplet formé du nom de la variable, de l'environnement de définition de la fonction et de son corps.

Les règles d'évaluation correspondantes, dites avec *liaison statique*, sont :

$$\begin{array}{c}
\overline{\Gamma \vdash \text{Fun}(x, e) \Rightarrow (x, \Gamma, e)} \\
\\
\frac{\Gamma \vdash e_1 \Rightarrow (x, \Gamma', e) \quad \Gamma \vdash e_2 \Rightarrow v_2 \quad \Gamma' \cdot \{x = v_2\} \vdash e \Rightarrow v_3}{\Gamma \vdash \text{App}(e_1, e_2) \Rightarrow v_3}
\end{array}$$

Dans l'exemple ci-dessus, l'évaluation de `(f 2)` retourne $(y, \{x = 2\}, 2*x+3*y)$ et donc l'évaluation de `((f 2) 3)` s'effectuera correctement.

Théorème 5.1 *Le langage des expressions arithmétiques fonctionnelles, avec la sémantique par liaison statique définie ci-dessus, est fortement typé.*

5.6 Codage des expressions *let*

Montrons maintenant le codage des expressions *let* étudiées au chapitre 4 que l'on trouve dans les langages fonctionnels habituels. Dans la mesure où ce codage exprime la signification des différentes expressions *let* en fonction des primitives d'abstraction et d'application, nous pouvons le décrire avec les règles sémantiques données figure 5.3.

La seconde règle se ramène en fait à la première, en considérant qu'une fonction d'un argument est un identificateur dont le corps est une abstraction. Ces règles peuvent surprendre, puisque c'est le corps de la fonction qui est passé en argument de l'expression g qui figure après le mot clé `in` dans la syntaxe concrète. Mais cela est assez naturel, puisque l'expression g utilise la fonction définie par le *let*, qui va donc être tout normalement substituée par son corps ...

La traduction que nous venons de présenter pour les *let fun* ne fonctionne que pour les définitions non récursives, une autre traduction est nécessaire pour les fonctions récursives, donc pour les expressions *let fun rec*. Cela explique la présence de deux mots clé

distincts, `let fun` et `let fun rec` pour définir d'une part les fonctions non récursives et d'autre part les fonctions récursives dans les langages de la famille ML.

Chapitre 6

Types structurés

Ce chapitre s'intéresse aux problèmes de sémantiques liés à l'utilisation de types appelés *structurés*, qui désignent d'ordinaire des types non fonctionnels obtenus à partir d'autres types, en particulier des types de base, à l'aide de constructeurs de types.

On classe les types structurés en deux catégories : les types produits et les types sommes. Les types produits correspondent à la construction d'un type structuré par produit cartésien de types déjà définis, un élément d'un type produit sera donc défini comme un n -uplets d'éléments ; alors que les types sommes correspondent à la construction d'un type structuré par union disjointe, un élément d'un type somme sera donc défini par une valeur parmi l'union des éléments possibles.

6.1 Les types produits

Nous nous intéressons d'abord aux types produits. Nous les classons en deux classes : les produits anonymes et les produits étiquetés. Les produits anonymes correspondent à la définition classique des n -uplets en mathématiques, alors que les produits étiquetés correspondent à la structure d'enregistrement que l'on trouve dans la plupart des langages de programmation.

Pour simplifier l'étude nous considérons les types couples, la généralisation aux n -uplets avec $n \geq 3$ étant immédiate.

6.1.1 Le type couple

On notera par $t \times t'$ le type des couples d'éléments de types respectifs t et t' . La grammaire abstraite des types est donc la suivante :

$$T \rightarrow \times (T, T) \mid \rightarrow (T, T) \mid Nat \mid Bool$$

Syntaxe

Donnons tout d'abord la syntaxe abstraite des expressions de type couple :

$$E \rightarrow Couple(E, E) \mid Fst(E) \mid Snd(E) \mid \dots$$

puis une syntaxe concrète avec les règles de calcul de la syntaxe abstraite :

$$\begin{array}{ll} E \rightarrow (E, E) & \{ Couple(E, E) \} \\ E \rightarrow fst(E) & \{ Fst(E) \} \\ E \rightarrow snd(E) & \{ Snd(E) \} \\ E \rightarrow \dots & \{ \dots \} \end{array}$$

Les points de suspension figurent les règles étudiées au cours des chapitres antérieurs, qui peuvent toutes être récupérées telles quelles.

Typage

Nous aurons besoin de trois nouvelles règles de typage, une pour chaque nouvelle forme d'expression abstraite :

$$\frac{\Gamma \vdash e_1 : t_1 \quad \Gamma \vdash e_2 : t_2}{\Gamma \vdash \text{Couple}(e_1, e_2) : \times (t_1, t_2)}$$

$$\frac{\Gamma \vdash e : t_1 \times t_2}{\text{Fst}(e) : t_1}$$

$$\frac{\Gamma \vdash e : t_1 \times t_2}{\text{Snd}(e) : t_2}$$

Évaluation

Nous avons de même trois nouvelles règles d'évaluation :

$$\frac{\Gamma \vdash e_1 \Rightarrow v_1 \quad \Gamma \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{Couple}(e_1, e_2) \Rightarrow (v_1, v_2)}$$

$$\frac{\Gamma \vdash e \Rightarrow (v_1, v_2)}{\text{Fst}(e) \Rightarrow v_1}$$

$$\frac{\Gamma \vdash e \Rightarrow (v_1, v_2)}{\text{Snd}(e) \Rightarrow v_2}$$

6.1.2 Les types enregistrements

Le type enregistrement permet de structurer des données de types variés que l'on peut ensuite retrouver via une clé d'accès. La grammaire des types devient la suivante:

$$T \rightarrow \text{TypeRecord}(TR) \mid \text{NomType}(String)$$

$$TR \rightarrow TRvide \mid TRcons(String, T, TR)$$

Par exemple, au type enregistrement (écrit en syntaxe concrète) $\{x:\text{int}, y:\text{int}\}$ correspond la syntaxe abstraite $TRcons("x", \text{Int}, TRcons("y", \text{Int}, TRvide))$.

Donnons maintenant la syntaxe abstraite des expressions de type enregistrement :

$$E \rightarrow \text{ExprRecord}(ER) \mid \text{AccesRecord}(E, String)$$

$$ER \rightarrow ERvide \mid TRcons(String, E, ER)$$

Par exemple, à l'expression (écrite en syntaxe concrète) $\{x=1, y=2\}$ correspond la syntaxe abstraite $ERcons("x", \text{Cte}(1), ERcons("y", \text{Cte}(2), ERvide))$

Pour des raisons qui seront explicitées dans le reste de ce paragraphe, un programme pourra posséder des déclarations de types. Nos programmes seront définis par la syntaxe concrète suivante, accompagnée de règles de calcul de la syntaxe abstraite :

$E \rightarrow \text{type } Id = T \text{ in } E$	$\{Typedecl(Id, T, E)\}$
$E \rightarrow E.Id$	$\{AccesRecord(E, \text{val}(Id_1))\}$
$E \rightarrow \{Id = E; ER$	$\{ExprREcord(ERcons(\text{val}(Id_1), E_1, ER_1))\}$
$ER \rightarrow \}$	$\{ERvide\}$
$ER \rightarrow Id = E, ER$	$\{ERcons(\text{val}(Id_1), E_1, ER_1)\}$
$T \rightarrow recordId : T, TR$	$\{TypeRecord(TRcons(\text{val}(Id_1), T_1, TR_1))\}$
$T \rightarrow Id$	$\{Nomtype(\text{val}(Id))\}$
$TR \rightarrow endrecord$	$\{TRvide\}$
$TR \rightarrow Id : T; TR$	$\{TRcons(\text{val}(Id_1), T_1, TR)\}$

Ainsi, au programme

```

type point = record
    x:int ;
    y:int
endrecord

in
let A : point = { x=3 ; y=4 }
in
let fun f(p:point):real =
    p.x * p.x + p.y * p.y
in f(A)

```

est associé la syntaxe abstraite

```

Typedecl(point, TypeRecord(TRcons(x, \Int, TRcons(y, \Int, TRvide))),
  Let(A, ExprRecord(ERcons(x, \Cte(3), ERcons(y, \Cte(4), ERvide))),
    Letfun(f, p, Nomtype(point), \Real,
      +(\times(AccesRecord(\Var(p), x), AccesRecord(\Var(p), x)),
        \times(AccesRecord(\Var(p), x), AccesRecord(\Var(p), x))),
      \App(f, \Var(A))))))

```

Typage

$$\frac{\Gamma, \Gamma_r \cdot \{x = t\} \vdash er : tr}{\Gamma, \Gamma_r \vdash Typedcl(x, tr, er)}$$

$$\frac{\Gamma, \Gamma_r \vdash er : tr}{\Gamma \vdash ExprRecord(er) : TypeRecord(tr)}$$

$$\frac{\Gamma \vdash E : TypeRecord(tr)}{\Gamma \vdash AccesRecord(E, x) : t}$$

6.2 Les types sommes

L'étude des types sommes est en quelque sorte duale de celle des types produits. Nous étudions en premier le cas de la somme anonyme de deux types, avant des regarder les sommes étiquetées quelconques.

6.2.1 Somme anonyme de deux types

On notera par $t \mid t'$ le type somme des types t et t' . La grammaire abstraite des types est donc la suivante :

$$T \rightarrow \text{Union}(T, T) \mid \rightarrow (T, T) \mid \text{Nat} \mid \text{Bool}$$

Syntaxe

Le type couple possède une fonction de construction pour former un couple à partir de deux éléments, et deux fonctions d'accès pour accéder soit au premier, soit au second élément du couple. Pour un type somme de deux types t_1 et t_2 , la situation est inversée : il y a deux fonctions de constructions que l'on notera arbitrairement *Premier* (respectivement *Second*) permettant de construire un élément du type somme $t_1 \mid t_2$ à partir d'un élément de type t_1 (respectivement t_2). La fonction d'accès correspondante se trouve être la construction bien connue de distinction de cas que l'on notera ici par **case** ... **of**. Voici par exemple une fonction qui prend en argument un entier ou un booléen et retourne le successeur si c'est un entier et 0 si c'est un booléen

```
let f = fun x : (Nat|Bool) ->
  case x of
    Premier(n) -> n+1
  | Second(b) -> 0
in
(f (Premier(1)))+(f (Second(true)))
```

La syntaxe abstraite des expressions de type somme est :

$$E \rightarrow \text{Premier}(E, E) \mid \text{Second}(E) \mid \text{Cas}(E, \text{Id}, E, \text{Id}, E) \mid \dots$$

Donnons une syntaxe concrète avec les règles de calcul de la syntaxe abstraite :

$$\begin{array}{ll} E \rightarrow \text{Premier}(E) & \{\text{Premier}(E_1)\} \\ E \rightarrow \text{Second}(E) & \{\text{Second}(E_1)\} \\ E \rightarrow \text{case } E \text{ of } \text{Premier}(\text{Id}) \rightarrow E \mid \text{Second}(\text{Id}) \rightarrow E \text{ endcase} & \{\text{Cas}(E_1, \text{val}(\text{Id}_1), E_2, \text{val}(\text{Id}_1), E_3)\} \\ E \rightarrow \dots & \{\dots\} \end{array}$$

Les points de suspension figurent les règles étudiées au cours des chapitres antérieurs, qui peuvent toutes être récupérées telles quelles.

Typage

Nous aurons besoin de trois nouvelles règles de typage, une pour chaque nouvelle forme d'expression abstraite :

$$\frac{\Gamma \vdash e_1 : t_1 \mid t_2 \quad \Gamma \vdash e_2 : t_3 \quad \Gamma \vdash e_3 : t_3}{\Gamma \vdash \text{Cas}(e_1, x_1, e_2, x_2, e_3) : t_3}$$

$$\frac{\Gamma \vdash e : t_1}{\text{Premier}(e) : t_1 \mid t_2}$$

$$\frac{\Gamma \vdash e : t_2}{\text{Second}(e) : t_1 \mid t_2}$$

Noter que la règle du **case** exprime bien que les deux expressions des deux alternatives doivent être de même type.

Évaluation

Nous avons maintenant quatre nouvelles règles d'évaluation :

$$\frac{\Gamma \vdash e_1 \Rightarrow \text{Premier}(v_1) \quad \Gamma \cdot \{x_1 = v_1\} \vdash e_2 \Rightarrow v_2}{\Gamma \vdash \text{Cas}(e_1, x_1, e_2, x_2, e_3) \Rightarrow v_2}$$

$$\frac{\Gamma \vdash e_1 \Rightarrow \text{Second}(v_1) \quad \Gamma \cdot \{x_2 = v_1\} \vdash e_3 \Rightarrow v_2}{\Gamma \vdash \text{Cas}(e_1, x_1, e_2, x_2, e_3) \Rightarrow v_2}$$

$$\frac{\Gamma \vdash e \Rightarrow v}{\text{Premier}(e) \Rightarrow \text{Premier}(v)}$$

$$\frac{\Gamma \vdash e \Rightarrow v}{\text{Second}(e) \Rightarrow \text{Second}(v)}$$

6.2.2 Sommes étiquetés

On trouve de tels types sous des dénominations variées dans les langages de programmation. En PASCAL, on trouve les enregistrements avec variantes qui permettent une définition de sommes. En langage C, il s'agit des types `union`. Dans les langages fonctionnels de la famille ML, il s'agit des types construits.

Chapitre 7

Sémantique d'un langage impératif

Ce chapitre s'intéresse à la sémantique des langages impératifs. Nous présentons un langage impératif avec affectations, instructions conditionnelles et boucles `while`. Nous définissons ensuite une traduction des programmes de ce langage impératif en des expressions fonctionnelles comme celles présentées dans le chapitre 4. Le but de cette traduction est de montrer comment les définitions formelles de sémantique d'un langage que nous avons données permettent de prouver de manière rigoureuse la correction d'une traduction d'un langage vers un autre.

7.1 Syntaxe concrète et abstraite

Notre langage impératif utilise seulement les types `bool` et `int`. Un programme commence avec une séquence (qui peut être vide) de déclarations des variables avec leurs valeurs initiales, suivi par une séquence d'instructions. Le seul type d'instruction simple est l'affectation; les instructions peuvent être enchaînées avec le symbole « ; », et peuvent être composées par des « parenthèses » `begin...end`, les boucles `while` et l'instruction conditionnelle. Voilà un exemple d'un programme pour calculer la valeur de 5^3 . Après l'exécution du programme, la variable `resultat` contiendra le résultat 125.

```
base:      int = 5;
exposant:  int = 3;
resultat:  int = base;
index:     int = 1;
while exposant > index do begin
  resultat := resultat * exposant;
  index := index+1;
end;
```

La syntaxe abstraite choisie pour notre langage impératif est donnée par la grammaire suivante. Nous utilisons pour les expressions arithmético-logiques la grammaire donnée dans la section 3.2 :

$$\begin{aligned}
P &\rightarrow \text{Prog}(D, I) \\
D &\rightarrow \text{Vide} \mid \text{Declare}(\text{String}, T, E, D) \\
I &\rightarrow \text{Skip} \mid \text{Cons}(I, I) \mid \text{Affect}(\text{String}, E) \mid \text{Cond}(E, I, I) \mid \text{While}(E, I) \\
T &\rightarrow \text{Bool} \mid \text{Int} \\
E &\rightarrow \dots \quad (\text{voir section 3.2})
\end{aligned}$$

$\frac{\Gamma \vdash i}{\Gamma \vdash \text{Prog}(\text{Vide}, i)}$	$\frac{\Gamma \cdot \{s : t\} \vdash \text{Prog}(d, i) \quad \Gamma \vdash e : t}{\Gamma \vdash \text{Prog}(\text{Declare}(s, t, e, d), i)} \quad \text{si } s \notin \Gamma$
$\frac{}{\Gamma \vdash \text{Skip}}$	$\frac{\Gamma \vdash i_1 \quad \Gamma \vdash i_2}{\Gamma \vdash \text{Cons}(i_1, i_2)}$
$\frac{\Gamma \vdash e : \text{Bool} \quad \Gamma \vdash i}{\Gamma \vdash \text{While}(e, i)}$	$\frac{\Gamma \vdash e : \text{Bool} \quad \Gamma \vdash i_1 \quad \Gamma \vdash i_2}{\Gamma \vdash \text{Cond}(e, i_1, i_2)}$
$\frac{\Gamma \vdash e : t}{\Gamma \vdash \text{Affect}(s, e)} \quad \text{si } s : t \in \Gamma$	

FIG. 7.1 – Typage des programmes impératifs.

L'axiome de cette grammaire est P .

Pour la grammaire concrète non-ambiguë et les règles de calcul de la syntaxe abstraite nous employons également les règles données dans section 3.2:

$P \rightarrow DS \ IS$	$\{\text{Prog}(DS_1, IS_1)\}$
$DS \rightarrow \varepsilon$	$\{\text{Vide}\}$
$DS \rightarrow \text{String} : T = E ; DS$	$\{\text{Declare}(\text{val}(\text{String}_1), T_1, E_1, DS_1)\}$
$IS \rightarrow \varepsilon$	$\{\text{Skip}\}$
$IS \rightarrow I ; IS$	$\{\text{Cons}(I_1, IS_1)\}$
$I \rightarrow \text{String} := E$	$\{\text{Affect}(\text{val}(\text{String}_1), E_1)\}$
$I \rightarrow \text{if } E \text{ then } I \text{ else } I$	$\{\text{Cond}(E_1, I_1, I_2)\}$
$I \rightarrow \text{while } E \text{ do } I$	$\{\text{While}(E_1, I_1)\}$
$I \rightarrow \text{begin } IS \text{ end}$	$\{IS_1\}$
$E \rightarrow \dots$	(voir section 3.2)

La syntaxe abstraite correspondante au programme en syntaxe concrète donné au-dessus est:

```

Prog(Declare(base, Int, 5, Declare(exposant, Int, 3,
  Declare(resultat, Int, Var(base), Declare(index, Int, 1, Vide))),
  Cons(While(> (Var(exposant), Var(index)),
    Cons(Affect(Var(resultat), ×(Var(resultat), Var(exposant))),
      Cons(Affect(Var(index), +(Var(index), Cte(1))),
        Skip))),
    Skip))

```

7.2 Typage

Les règles de typage sont données sur la figure 7.1. On utilise deux sortes de jugements :

1. des jugements $\Gamma \vdash e$ où e est une expression de la syntaxe abstraite de sorte P ou I , disant que e est bien typé dans l'environnement Γ ;
2. des jugements $\Gamma \vdash e : t$ où e est une expression de la syntaxe abstraite de la sorte E et t est un type, disant que e a dans l'environnement Γ le type t . Les règles pour les jugements de ce type sont données sur la figure 3.6.

$\frac{\Gamma \vdash i \Rightarrow \Gamma'}{\Gamma \vdash \text{Prog}(\text{Vide}, i) \Rightarrow \Gamma'}$	$\frac{\Gamma \vdash e \Rightarrow v \quad \Gamma \cdot \{s = v\} \vdash \text{Prog}(d, i) \Rightarrow \Gamma'}{\Gamma \vdash \text{Prog}(\text{Declare}(s, t, e, d), i) \Rightarrow \Gamma'}$
$\frac{}{\Gamma \vdash \text{Skip} \Rightarrow \Gamma}$	$\frac{\Gamma \vdash i_1 \Rightarrow \Gamma' \quad \Gamma' \vdash i_2 \Rightarrow \Gamma''}{\Gamma \vdash \text{Cons}(i_1, i_2) \Rightarrow \Gamma''}$
$\frac{\Gamma \vdash e \Rightarrow \text{true} \quad \Gamma \vdash i_1 \Rightarrow \Gamma'}{\Gamma \vdash \text{Cond}(e, i_1, i_2) \Rightarrow \Gamma'}$	$\frac{\Gamma \vdash e \Rightarrow \text{false} \quad \Gamma \vdash i_2 \Rightarrow \Gamma'}{\Gamma \vdash \text{Cond}(e, i_1, i_2) \Rightarrow \Gamma'}$
$\frac{\Gamma \vdash e \Rightarrow v}{\Gamma \vdash \text{Affect}(s, e) \Rightarrow \Gamma \cdot \{s = v\}}$	$\frac{\Gamma \vdash e \Rightarrow \text{false}}{\Gamma \vdash \text{While}(e, i) \Rightarrow \Gamma'}$
$\frac{\Gamma \vdash e \Rightarrow \text{true} \quad \Gamma \vdash i \Rightarrow \Gamma' \quad \Gamma' \vdash \text{While}(e, i) \Rightarrow \Gamma''}{\Gamma \vdash \text{While}(e, i) \Rightarrow \Gamma''}$	

FIG. 7.2 – Règles d'évaluations des programmes impératifs.

7.3 Évaluation

Au contraire des programmes fonctionnels définis dans le chapitre 4, les programmes de notre langage impératif ne disposent pas d'une notation pour leur résultat. Afin de définir un langage de programmation pratique il faudrait bien sûr ajouter, parmi d'autres, des instructions de sorties. Pour raison de simplicité nous nous contraindrons ici à un langage minimal qui nous permet d'étudier les constructions essentielles de la programmation impérative. Le résultat de l'exécution d'un programme impératif sera pour l'ensemble des valeurs de ses variables.

Nous définissons la sémantique des programmes impératifs à l'aide des jugements de la forme $\Gamma \vdash i \Rightarrow \Gamma'$, disant que l'exécution de l'instruction (ou du programme) i dans l'environnement Γ termine et produit l'environnement Γ' . Les règles d'évaluation sont données sur la figure 7.2. On utilise, comme pour les règles de typage, des jugements d'évaluation des expressions arithmético-logiques définis sur la figure 3.7.

7.4 Traduction des programmes impératifs en programmes fonctionnels récursifs

Nous allons maintenant traduire les programmes impératifs en des programmes fonctionnels comme définis dans la section 4.2, *en conservant la sémantique des programmes*. Le problème est que les sémantiques des deux langages sont définies par des formalismes différents : la sémantique d'un programme fonctionnel est définie comme une fonction qui renvoie, pour un environnement donné, une valeur. La sémantique d'un programme impératif, par contre, est une fonction qui renvoie un environnement. Parce que l'exécution d'un programme impératif donné peut seulement changer les valeurs des variables qui se trouvent dans le programme, (par exemple des variables $x_1, \dots, x_n$), on peut préciser l'objectif de la traduction comme suit : on traduit un programme impératif p avec les variables $x_1, \dots, x_n$ de types respectifs $t_1, \dots, t_n$ en une expression fonctionnelle f du type $(t_1, \dots, t_n)$ tel que si $v_1, \dots, v_n$ sont les valeurs des variables $x_1, \dots, x_n$ après l'exécution de p , alors l'évaluation de f renvoie le n -uplet $(v_1, \dots, v_n)$.

Afin de ne pas alourdir la traduction par des détails, nous sommes obligés de nous contraindre ici au cas des programmes qui utilisent une seule variable. Nous supposons doré-

$\frac{i \Rightarrow f}{\text{Prog}(\text{Vide}, i) \Rightarrow f}$	$\frac{\text{Prog}(d, i) \Rightarrow f}{\text{Prog}(\text{Declare}(x, \text{Int}, e, d), i) \Rightarrow \text{Let}(x, e, f)}$
$\frac{}{\text{Skip} \Rightarrow x}$	$\frac{i_1 \Rightarrow f_1 \quad i_2 \Rightarrow f_2}{\text{Cond}(e, i_1, i_2) \Rightarrow \text{If}(e, f_1, f_2)}$
$\frac{i_1 \Rightarrow f_1 \quad i_2 \Rightarrow f_2}{\text{Cons}(i_1, i_2) \Rightarrow \text{Let}(x, f_1, f_2)}$	$\frac{}{\text{Affect}(x, e) \Rightarrow e}$
$\frac{i \Rightarrow f}{\text{While}(e, i) \Rightarrow \text{LetFunRec}(l, x, \text{Int}, \text{Int}, \text{If}(e, \text{App}(l, f), \text{Var}(x)), \text{App}(l, (\text{Var}(x))))}$	

FIG. 7.3 – Traduction des programmes impératifs, avec une seule variable x du type `int`, aux programmes fonctionnels.

navant que cette variable s'appelle x , et qu'elle a le type `int`. On pourrait généraliser cette traduction au cas des programmes impératifs utilisant un nombre quelconque de variables à l'aide des types structurés présentés dans le chapitre 6 (voir exercice 7.4).

On obtient, par exemple, pour le programme

```
x: int = 7;
while x > 1 do begin
  if 2*(x/2) = x
  then x := x/2
  else x := 3*x+1
end
```

le programme fonctionnel suivant:

```
let x = 7
in let rec l(x:int):int =
  if x > 1
  then l(if 2*(x/2)=x then x/2 else 3*x+1)
  else x
in l(x)
```

Notons que la boucle `while` d'un programme impératif se traduit en une fonction récursive.

Les règles pour les jugements de la forme $i \Rightarrow f$, exprimant que le programme fonctionnel f est la traduction du programme impératif i , sont données par la figure 7.3.

Nous supposons que les identificateurs x et l n'apparaissent pas dans i . Notons que les expressions arithmético-logiques font partie du langage impératif et du langage fonctionnel, et que la traduction est par conséquent immédiate.

Théorème 7.1 *Soit p un programme impératif avec pour seul identificateur x . Soit f l'expression fonctionnelle traduction de p , c'est-à-dire telle que le jugement $p \Rightarrow f$ soit valide. Soit Γ, Γ' des environnements quelconques. On a la proposition suivante :*

$$\text{Si} \quad \Gamma \vdash p \Rightarrow \Gamma' \quad \text{alors} \quad \Gamma \vdash f \Rightarrow v' \quad \text{avec} \quad x = v' \in \Gamma'$$

Autrement dit, la valeur de f est exactement la valeur de x dans l'environnement obtenu par exécution de p , ce qui est bien ce que l'on souhaite pour dire que la traduction est correcte.

Comment peut-on aborder la preuve du théorème 7.1? Une première idée est une preuve par récurrence sur la structure du programme impératif i . Regardons les règles de la figure 7.2, nous remarquons que pour toutes les constructions *sauf le While* les programmes dans les hypothèses des règles sont des sous-expressions du programme de la conclusion. Par conséquent, une récurrence sur la *taille* du programme, c'est-à-dire sur le nombre d'étiquettes abstraites qu'il contient, peut réussir dans tous les cas sauf le While. Pour le While, le programme de la conclusion se trouve également dans l'hypothèse dans le cas où la condition est satisfaite. Dans ce cas, nous sommes donc obligés d'utiliser une technique appropriée à la nature de l'instruction While, c'est une récurrence sur le nombre d'itérations de la boucle. Ces deux idées sont réunies dans une *récurrence sur la taille de l'arbre de dérivation* du jugement $\Gamma \vdash i \Rightarrow \Gamma'$.

Preuve:

Avant de prouver le théorème 7.1 nous montrons une propriété analogue pour les instructions :

Quelle que soit l'instruction i avec pour seul identificateur x , si $\Gamma \vdash i \Rightarrow \Gamma'$ avec $x \in \Gamma$ et $i \Rightarrow f$, alors $\Gamma \vdash f \Rightarrow v'$ avec $\{x = v'\} \in \Gamma'$.

Nous commençons avec les deux cas des instructions simples :

Cas : $i = \text{Skip}$

On a $\Gamma \vdash \text{Skip} \Rightarrow \Gamma'$, donc il faut par les règles d'évaluation des programmes impératifs que $\Gamma' = \Gamma$. Selon les règles de traduction, $\text{Skip} \Rightarrow \text{Var}(x)$.

Donc, il faut montrer que $\Gamma \vdash \text{Var}(x) \Rightarrow v$ où $\{x = v\} \in \Gamma$. C'est une conséquence triviale des règles d'évaluation des expressions arithmético-logiques et du fait que $\{x = v\} \in \Gamma$.

Cas : $i = \text{Affect}(x, e)$

On a $\Gamma \vdash \text{Affect}(x, e) \Rightarrow \Gamma'$, donc il faut par les règles d'évaluation des programmes impératifs qu'il y a un v tel que $\Gamma \vdash e \Rightarrow v$ et $\Gamma' = \Gamma \cdot \{x = v\}$. Selon les règles de traduction, $\text{Affect}(x, e) \Rightarrow e$.

Donc il faut montrer que $\Gamma \vdash e \Rightarrow v'$ où $\{x = v'\} \in \Gamma \cdot \{x = v\}$. C'est une conséquence triviale de l'hypothèse $\Gamma \vdash e \Rightarrow v$ et du fait que $v' = v$.

Dans les trois cas des instructions composées nous sommes amenés à faire une preuve par récurrence.

Cas : $i = \text{Cond}(e, j_1, j_2)$

On a $\Gamma \vdash \text{Cond}(e, j_1, j_2) \Rightarrow \Gamma'$. Selon les règles de traduction, il existe f_1 et f_2 tels que

$$\begin{aligned} j_1 &\Rightarrow f_1 \\ j_2 &\Rightarrow f_2 \\ \text{Cond}(e, j_1, j_2) &\Rightarrow \text{If}(e, f_1, f_2) \end{aligned}$$

Il faut montrer que $\Gamma \vdash \text{If}(e, f_1, f_2) \Rightarrow v'$ où $\{x = v'\} \in \Gamma'$.

Sous-cas : $\Gamma \vdash e \Rightarrow \text{true}$

Par les règles d'évaluation des programmes impératifs on obtient

$$\Gamma \vdash j_1 \Rightarrow \Gamma'$$

L'application de l'hypothèse de récurrence à j_1 donne

$$\Gamma \vdash f_1 \Rightarrow v'$$

Donc, on obtient par les règles d'évaluation des programmes fonctionnels :

$$\frac{\frac{\vdots}{\Gamma \vdash e \Rightarrow \mathbf{true}} \quad \frac{\vdots}{\Gamma \vdash f_1 \Rightarrow v'}}{\Gamma \vdash \text{If}(e, f_1, f_2) \Rightarrow v'}$$

Sous-cas : $\Gamma \vdash e \Rightarrow \mathbf{false}$

Ce cas est analogue au cas précédent.

Cas : $i = \text{Cons}(j_1, j_2)$

On a $\Gamma \vdash \text{Cons}(j_1, j_2) \Rightarrow \Gamma'$. Selon les règles de traduction il existe f_1 et f_2 tels que

$$\begin{aligned} j_1 &\Rightarrow f_1 \\ j_2 &\Rightarrow f_2 \\ \text{Cons}(j_1, j_2) &\Rightarrow \text{Let}(x, f_1, f_2) \end{aligned}$$

Il faut montrer que $\Gamma \vdash \text{Let}(x, f_1, f_2) \Rightarrow v'$ où $\{x = v'\} \in \Gamma'$. Selon les règles d'évaluation des programmes impératifs il existe Γ'' tel que

$$\begin{aligned} \Gamma \vdash j_1 &\Rightarrow \Gamma'' \\ \Gamma'' \vdash j_2 &\Rightarrow \Gamma' \end{aligned}$$

Soit $\{x = v''\} \in \Gamma''$. Puisque x est le seul identificateur de j_2 , on obtient également

$$\Gamma \cdot \{x = v''\} \vdash j_2 \Rightarrow \Gamma'$$

L'application de l'hypothèse de récurrence à j_1 et j_2 donne

$$\begin{aligned} \Gamma \vdash f_1 &\Rightarrow v'' \\ \Gamma \cdot \{x = v''\} \vdash f_2 &\Rightarrow v' \end{aligned}$$

Donc, on obtient par les règles d'évaluation des programmes fonctionnels

$$\frac{\frac{\vdots}{\Gamma \vdash f_1 \Rightarrow v''} \quad \frac{\vdots}{\Gamma \cdot \{x = v''\} \vdash f_2 \Rightarrow v'}}{\Gamma \vdash \text{Let}(x, f_1, f_2) \Rightarrow v'}$$

Cas : $i = \text{While}(e, j)$

On a $\Gamma \vdash \text{While}(e, j) \Rightarrow \Gamma'$. Selon les règles de traduction il existe f tel que

$$\begin{aligned} j &\Rightarrow f \\ \text{While}(e, j) &\Rightarrow \text{LetFunRec}(l, x, \text{Int}, \text{Int}, \text{If}(e, \text{App}(l, f), \text{Var}(x)), \text{App}(l, (\text{Var}(x)))) \end{aligned}$$

Il faut montrer que

$$\Gamma \vdash \text{LetFunRec}(l, x, \text{Int}, \text{Int}, \text{If}(e, \text{App}(l, f), \text{Var}(x)), \text{App}(l, (\text{Var}(x)))) \Rightarrow v'$$

où $x = v' \in \Gamma'$. Soit $x = v \in \Gamma$ et $\Gamma_1 = \Gamma \cdot \{l = (x, \text{If}(e, \text{App}(l, f), \text{Var}(x)))\}$.

Sous-cas : $\Gamma \vdash e \Rightarrow \mathbf{false}$

Puisque l n'apparaît pas dans e et puisque $\{x = v\} \in \Gamma$ on a aussi

$$\Gamma_1 \cdot \{x = v\} \vdash e \Rightarrow \mathbf{false}$$

Par les règles de sémantique impérative il faut que $\Gamma = \Gamma'$, et par conséquent que $v = v'$.
On construit la dérivation suivante:

$$\frac{\Gamma_1 \vdash \text{Var}(x) \Rightarrow v \quad \frac{\frac{\vdots}{\Gamma_1 \cdot \{x = v\} \vdash e \Rightarrow \text{false}} \quad \Gamma_1 \cdot \{x = v\} \vdash \text{Var}(x) \Rightarrow v}{\Gamma_1 \cdot \{x = v\} \vdash \text{If}(e, \text{App}(l, f), \text{Var}(x)) \Rightarrow v}}{\Gamma_1 \vdash \text{App}(l, (\text{Var}(x))) \Rightarrow v}}{\Gamma \vdash \text{LetFunRec}(l, x, \text{Int}, \text{Int}, \text{If}(e, \text{App}(l, f), \text{Var}(x)), \text{App}(l, (\text{Var}(x)))) \Rightarrow v}$$

Sous-cas : $\Gamma \vdash e \Rightarrow \text{true}$

Par les règles d'évaluation des programmes impératifs, il existe Γ'' tel que

$$\begin{aligned} \Gamma \vdash i &\Rightarrow \Gamma'' \\ \Gamma'' \vdash \text{While}(e, i) &\Rightarrow \Gamma' \end{aligned}$$

Puisque x est le seul identificateur dans $\text{While}(e, i)$ on obtient également

$$\Gamma_1 \cdot \{x = v\} \vdash \text{While}(e, i) \Rightarrow \Gamma'$$

L'application de l'hypothèse de récurrence donne

$$\Gamma_1 \cdot \{x = v\} \vdash \text{LetFunRec}(l, x, \text{Int}, \text{Int}, \text{If}(e, \text{App}(l, f), \text{Var}(x)), \text{App}(l, (\text{Var}(x)))) \Rightarrow v'$$

Donc, selon les règles d'évaluation des programmes fonctionnels et puisque Γ_1 contient déjà la définition de l :

$$\Gamma_1 \cdot \{x = v\} \vdash \text{App}(l, (\text{Var}(x))) \Rightarrow v'$$

On construit donc la dérivation suivante:

$$\frac{\Gamma_1 \vdash \text{Var}(x) \Rightarrow v \quad \frac{\frac{\vdots}{\Gamma_1 \cdot \{x = v\} \vdash e \Rightarrow \text{true}} \quad \frac{\vdots}{\Gamma_1 \cdot \{x = v\} \vdash \text{App}(l, f) \Rightarrow v'}}{\Gamma_1 \cdot \{x = v\} \vdash \text{If}(e, \text{App}(l, f), \text{Var}(x)) \Rightarrow v'}}{\Gamma_1 \vdash \text{App}(l, (\text{Var}(x))) \Rightarrow v'}}{\Gamma \vdash \text{LetFunRec}(l, x, \text{Int}, \text{Int}, \text{If}(e, \text{App}(l, f), \text{Var}(x)), \text{App}(l, (\text{Var}(x)))) \Rightarrow v'}$$

Afin de compléter la preuve de théorème 7.1, il nous reste à traiter les deux cas des programmes de la forme $i = \text{Prog}(\text{Vide}, j)$ et $i = \text{Prog}(\text{Declare}(s, t, e, d), i)$. Ces deux preuves sont très similaires aux cas traités au-dessus, voir exercice 7.3. □

7.5 Exercices

Exercice 7.1

On définit un deuxième système de règles pour l'évaluation des programmes impératifs, en remplaçant dans le système de la figure 7.2 les deux règles pour `While` par la règle

$$\frac{\Gamma \vdash \text{Cond}(e, \text{Cons}(i, \text{While}(e, i)), \text{Skip}) \Rightarrow \Gamma'}{\Gamma \vdash \text{While}(e, i) \Rightarrow \Gamma'}$$

Montrer que les deux systèmes de règles définissent la même sémantique.

Exercice 7.2

Étendre la syntaxe concrète, la syntaxe abstraite et les règles de typage et d'évaluation par les boucles `repeat...until`. Donner une règle de traduction en programmes fonctionnelles.

Exercice 7.3

Compléter la preuve de théorème 7.1 par les preuves des deux cas $i = \text{Prog}(\text{Vide}, j)$ et $i = \text{Prog}(\text{Declare}(s, t, e, d), i)$.

Correction :

Cas: $i = \text{Prog}(\text{Vide}, j)$

Par notre hypothèse, $\text{Prog}(\text{Vide}, j) \Rightarrow f$ est vrai. La seule règle qui permet de déduire un jugement de cette forme est la première règle de figure 7.3. Si on sait que $\text{Prog}(\text{Vide}, j) \Rightarrow f$ est vrai et qu'on a appliqué la première règle pour le déduire, il faut que son hypothèse est aussi vrai. Donc, $f \Rightarrow j$ est vrai.

On applique un raisonnement similaire aux jugements d'évaluation: Par hypothèse, $\Gamma \vdash \text{Prog}(\text{Vide}, j) \Rightarrow \Gamma'$ est vrai. La seule règle qui permet de déduire un tel jugement est la première règle de la figure 7.2. Donc, sa hypothèse $\Gamma \vdash j \Rightarrow \Gamma'$ est aussi vrai.

On sait donc que $f \Rightarrow j$ et $\Gamma \vdash j \Rightarrow \Gamma'$ sont vrai. La longueur de la dérivation de $f \Rightarrow j$ doit être inférieure à la longueur de la dérivation de $\text{Prog}(\text{Vide}, j) \Rightarrow f$, donc nous appliquons la hypothèse d'induction et obtenons que $\Gamma \vdash f \Rightarrow \Gamma'(x)$ est vrai.

Cas: $i = \text{Prog}(\text{Declare}(s, t, e, d), i)$

Par les règles de sémantique impérative, il existe v tel que $\Gamma \vdash e \Rightarrow v$ et $\Gamma \cdot \{s = v\} \vdash \text{Prog}(d, i) \Rightarrow \Gamma'$. Par les règles de traduction, $f = \text{Let}(s, e, g)$ où $\text{Prog}(d, i) \Rightarrow g$. Nous obtenons par l'hypothèse d'induction $\Gamma \cdot \{s = v\} \vdash g \Rightarrow \Gamma'(x)$ et, par les règles de sémantique fonctionnelle, $\Gamma \vdash f \Rightarrow \Gamma'(x)$.

Exercice 7.4

Généraliser la traduction du chapitre 7.4 au cas de programmes avec un nombre quelconque de variables.

=====

Troisième partie

Études de cas

Chapitre 1

Généralités sur les études de cas

1.1 Le contenu des rapports

1.1.1 Manuel d'utilisation

Cette partie s'adresse aux utilisateurs de votre programme. Il faut qu'un utilisateur trouve dans cette partie toutes les informations nécessaires pour utiliser votre programme. En particulier, vous expliquerez dans cette partie :

1. les fonctionnalités de votre programme : brièvement, sans répéter l'énoncé, mais en expliquant tous les ajouts, choix ou options que vous avez éventuellement été amenés à faire ;
2. comment lancer votre programme ;
3. la syntaxe que votre programme accepte. **Cette syntaxe sera donnée en deux parties, une pour les textes HTML, une pour les expressions et les définitions.** Pour chacune des deux grammaires vous définirez les unités lexicales par des tokens associés à des expressions rationnelles, puis la syntaxe complète par une grammaire sur le vocabulaire formé par les tokens ;
4. les cas d'erreur possibles et le comportement que devra avoir votre programme dans ces cas d'erreur.

1.1.2 Structure du programme et découpage en modules

Dans une seconde partie de votre rapport, vous décrirez le raisonnement que vous avez été amené à faire pour l'analyse descendante du programme, et de quelle façon vous en avez déduit le découpage en modules qui vous paraissait le mieux adapté. Vous dessinerez le graphe de dépendances entre les modules.

1.1.3 Description des modules

Dans une troisième partie du rapport, et pour chacun des modules, dans un ordre compatible avec le graphe de dépendance, vous donnerez dans un premier temps les spécifications de ce module :

1. son rôle en général ;

2. les types qui y sont définis, décrits abstraitement, c.-à-d. par leur noms et les spécifications des fonctions qui permettent de créer des valeurs de ce type, et de les manipuler. **vous décrierez les types de données pour les documents HTML, les expressions et les définitions par des grammaires abstraites ;**
3. les spécifications des fonctions qui y sont définies.

(Vous ne décrierez que les types et fonctions de l'interface du module). Vous décrierez dans un second temps la réalisation de ce module :

1. les structures de données que vous avez choisies pour représenter les types, et les raisons de ces choix ;
2. les algorithmes que vous avez choisis pour les fonctions, *lorsque ceux-ci ne sont pas triviaux.*

Si vous avez été amenés à définir des types privés ou des fonctions locales au module, vous les spécifierez.

Les algorithmes seront décrits à un haut niveau. Les fonctions d'analyse syntaxique devront être décrites par des grammaires avec actions. **L'évaluation des cases des tableaux sera décrite par des règles sémantiques s'appliquant sur les jugements appropriés. Chaque type de jugement devra être clairement défini.**

1.1.4 Jeu de tests

Dans une quatrième partie, vous donnerez enfin un ensemble de fichiers de tests que vous aurez choisis pour tester toutes les fonctionnalités de votre programme, sans oublier les cas d'erreurs. Vous donnerez les traces d'exécution de vos fichiers de tests, et vous commenterez les résultats obtenus. Vous expliquerez en particulier les cas où les tests ne se passent pas comme prévu, en essayant de trouver pourquoi. Les jeux de tests sont notés sur l'exhaustivité (c.-à-d. si les tests couvrent tous les cas d'exécution du programme), et aussi sur les commentaires expliquant pourquoi ces tests ont été choisis.

Enfin, il est inutile de rendre un listing de votre programme, mais par contre vous fournirez une disquette avec les fichiers sources (.ml et .mli) et les fichiers de tests. **Mettez vos noms et le nom de l'enseignant de votre groupe sur la disquette.**

1.1.5 Soutenance

Une soutenance se déroulera lors de votre séance de TP entre le mercredi 3 juin et le mardi 9 juin. Vous préparerez une petite présentation des fonctionnalités de votre programme sur des exemples bien choisis (environ 5mn), puis vous répondrez **individuellement** aux questions qui vont être posées.

Chapitre 2

Compilation d'un mini-Pascal

Chapitre 3

Un langage d'expressions rationnelles

Chapitre 4

Un langage pour la composition de documents

Dans cette étude de cas, nous allons écrire un programme de *composition* de documents, inspiré d'un logiciel déjà existant appelé LaTeX. Le but est de fournir en entrée une description de *haut niveau* d'un texte, et ce logiciel devra *mettre en page* le texte. Une description de haut niveau d'un texte consiste à décrire ces composants logiques : titre du document, titres et sous-titres de sections, énumérations, etc.; et la mise en page consiste alors à utiliser les caractères d'une fonte et d'une taille appropriée, couper les phrases et répartir les espaces entre les mots de manière à produire un texte avec des marges données, numérotter les titres et sous-titres de sections, numérotter les pages, choisir de changer de page aux endroits opportuns, etc.

4.1 Documents simples

Un fichier source type que l'on désire donner en entrée du programme est donné figure 4.1. Une sortie possible est alors donnée figure 4.2. Cette mise en page a été effectuée par un certain nombre de choix de style :

- la largeur du texte est de 80mm ;
- les paragraphes de texte doivent être « justifiés » (c.-à-d. marges alignées comme expliqué ci-dessus), dans la fonte Times de taille 3.5mm, le début de paragraphe devant être indenté de 7mm ;
- le titre du document doit être composé dans la fonte Helvetica, grasse, de taille 6mm ;
- un titre de section doit être composé dans la fonte Times, grasse, de taille 5mm ;
- un titre de sous-section doit être composé dans la fonte Times, oblique, de taille 4.2mm ;
- l'intervalle entre les lignes est de 1,2 fois la taille de la fonte ;
- etc.

Dans un premier temps, ces choix de style de composition pourront être codés directement dans le programme de composition. Nous proposerons à la section 4.1.5 une méthode pour paramétrer facilement le style.

Remarquer que les espaces entre les mots et les passages à la ligne du texte source n'ont aucune influence sur le texte de sortie : en effet les passages à la ligne et l'espacement entre

```
\title{Ceci est le titre du document}

Voici une première phrase introductive au document.

\section{Ceci est un titre de section}

Ceci est le premier paragraphe de la première section. Il contient
trois phrases. Noter la façon dont les espaces entre les mots sont
ajustés de manière à bien marquer les marges à droite et à
gauche.

Voici le deuxième paragraphe. Le changement de paragraphe est
matérialisé par une ligne vide.

\section{Voici la deuxième section}

\subsection{Et ceci est une sous-section}

Les sous-sections sont des sections de niveau deux, numérotées avec
deux nombres.

\subsection{Deuxième sous-section}

La numérotation des sections et sous-sections doit évidemment être
calculée automatiquement par le logiciel de composition.
```

FIG. 4.1 – *Exemple de fichier source simple*

les mots est calculé de manière à ce que les lignes soient contre les marges gauches et droite, en remplissant les lignes au maximum.

Il nous reste à parler de la façon de produire le texte de sortie : pour cela nous allons utiliser un autre langage appelé Postscript, qui est un langage de description de textes à imprimer. Il s'agit en fait du langage de commande pour certaines imprimantes. Comme il n'est pas question dans cette étude d'apprendre ce langage, nous fournissons un module PASCAL qui permet de produire des fichiers Postscript de manière simple. Ce module est décrit dans la section suivante.

4.1.3 Le module POSTSCR, et l'outil GhostView

Ce module fournit l'essentiel des procédures permettant de créer un fichier Postscript :

- définition des fontes possibles ;
- ouverture et fermeture d'un fichier Postscript ;
- écriture d'une chaîne de caractères dans une fonte donnée à une position donnée de la page ;
- détermination de la largeur d'une chaîne de caractères dans une fonte donnée (indispensable pour la mise en page) ;
- passage à la page suivante ;

L'interface PASCAL de ce module est donnée en annexe.

Ceci est le titre du document

Voici une première phrase introductive au document.

4.1 Ceci est un titre de section

Ceci est le premier paragraphe de la première section. Il contient trois phrases. Noter la façon dont les espaces entre les mots sont ajustés de manière à bien marquer les marges à droite et à gauche.

Voici le deuxième paragraphe. Le changement de paragraphe est matérialisé par une ligne vide.

4.2 Voici la deuxième section

4.2.1 *Et ceci est une sous-section*

Les sous-sections sont des sections de niveau deux, numérotées avec deux nombres.

4.2.2 *Deuxième sous-section*

La numérotation des sections et sous-sections doit évidemment être calculée automatiquement par le logiciel de composition.

FIG. 4.2 – Une page de sortie possible pour l'exemple

Ce module permet donc facilement de produire des fichiers Postscript. Ensuite, nous souhaitons bien sûr pouvoir visualiser les pages produites à l'écran, et aussi imprimer ces pages. Pour cela, nous disposons de l'outil Ghostview. Il se lance sous Windows en cliquant sur l'icône GSView. Dans le menu qui vous est proposé, vous pouvez ouvrir un fichier Postscript (File/Open), et (après quelques secondes d'attente...) la première page est visualisée. La barre d'outil sur le côté gauche permet de passer à la page suivante, de zoomer, etc. Dans les menus est également proposé d'imprimer la page visualisée (File/Print) ou tout le fichier (File/Print File).

4.1.4 Les fichiers sources et leur syntaxe

Pour résumer, du point de vue pratique, le programme de composition devra opérer de la manière suivante :

1. analyse syntaxique du fichier d'entrée ;

2. le cas échéant, signaler une erreur de syntaxe et s'arrêter ;
3. si aucune erreur de syntaxe n'est détectée, construction d'une structure de données mémorisant le texte source ;
4. application d'algorithmes de composition sur la structure de données produite, la sortie étant un fichier Postscript construit à l'aide du module fourni.

Comme on a pu le remarquer sur l'exemple, il existe certains mots clés qui sont précédés par un caractère `\` (backslash en anglais). En effet, les mots clés ou identificateurs que l'on trouve dans les langages de programmation classiques (`begin`, `end`, `if`, etc.) ne peuvent pas être utilisés ici car ce sont des mots que l'on peut utiliser dans le texte. Aussi, dans cette étude, les mots-clés et identificateurs seront toujours formés d'un `\` suivi d'une suite de lettres minuscules ou majuscules, ce qui donne ainsi `\title`, `\section` et `\subsection`. D'autre part, les parenthèses étant elles-mêmes des symboles de ponctuation habituels dans un texte, elles sont remplacées par les accolades.

La syntaxe des fichiers d'entrée doit suivre les règles suivantes :

- un texte est constitué d'une suite de *paragraphes* séparés par des lignes vides ;
- un paragraphe peut être soit un paragraphe *normal* formé par une séquence de *mots*, soit un paragraphe *spécial* (`title`, `section`, `subsection`) introduit par un mot clé puis une séquence de mots entre accolades.
- un mot est n'importe quelle suite de caractères non spéciaux (c.à-d. ni `\`, ni `{`, ni `}`, ni espace, tabulation ou retour à la ligne).

Comme montré sur l'exemple, les passages à la ligne du texte source n'ont pas d'influence sur les passages à la ligne du texte résultat. Aussi les espaces, tabulations et retour à la ligne pourront être considérés comme équivalents, et servent de séparateurs entre unités lexicales. Par contre, deux (ou plus) passages à la ligne successifs (éventuellement séparés par des espaces et tabulations) seront significatifs puisqu'ils séparent deux paragraphes : une telle séquence de caractères devra donc former une unité lexicale.

4.1.5 Extensions du langage

Le langage décrit jusqu'à présent permet de composer des documents très simples. Des constructions supplémentaires sont utiles pour des documents plus complexes.

Nous proposons ici quelques extensions, dans l'ordre de difficulté, et qui compteront pour quelques points dans la notation du rapport numéro 2.

Numérotation des pages et entêtes

Numéroter chaque page par un numéro placé en bas de page et centré. Introduire deux types de paragraphes supplémentaires :

```
\lefthead{texte}
\righthead{texte}
```

qui définissent des portions de texte à écrire sur chaque page, en haut à gauche pour `\lefthead`, en haut à droite pour `\righthead`.

α	'a'	\alpha	β	'b'	\beta	γ	'g'	\gamma
δ	'd'	\delta	ε	'e'	\epsilon	ζ	'z'	\zeta
η	'h'	\eta	θ	'q'	\theta	ι	'i'	\iota
κ	'k'	\kappa	λ	'l'	\lambda	μ	'm'	\mu
ν	'n'	\nu	ξ	'x'	\xi	$\omicron$	'o'	\omicron
π	'p'	\pi	ρ	'r'	\rho	σ	's'	\sigma
τ	't'	\tau	υ	'u'	\upsilon	ϕ	'j'	\phi
χ	'c'	\chi	ψ	'y'	\psi	ω	'w'	\omega
Γ	'G'	\Gamma	Δ	'D'	\Delta	Θ	'Q'	\Theta
Λ	'L'	\Lambda	Ξ	'X'	\Xi	Π	'P'	\Pi
Σ	'S'	\Sigma	Υ	'U'	\Upsilon	Φ	'F'	\Phi
Ψ	'Y'	\Psi	Ω	'W'	\Omega			
$\geq$	179	\geq	$\leq$	163	\leq	$\neq$	185	\neq
$\pm$	177	\pm	$\times$	180	\times	$\div$	184	\div
$\sum$	229	\sum	$\int$	242	\int	$\prod$	213	\prod
∞	165	\infty	$\bullet$	183	\bullet	$\aleph$	192	\aleph
$\emptyset$	198	\emptyset	$\in$	206	\in	$\subset$	204	\subset
$\subseteq$	205	\subseteq	$\cup$	200	\cup	$\cap$	199	\cap
$\leftarrow$	172	\leftarrow	$\rightarrow$	174	\rightarrow	$\uparrow$	173	\uparrow
$\downarrow$	175	\downarrow	$\leftrightarrow$	171	\leftrightarrow	$\Uparrow$	221	\Uparrow
$\Leftarrow$	220	\Leftarrow	$\Rightarrow$	222	\Rightarrow			
$\Downarrow$	223	\Downarrow	$\Leftrightarrow$	219	\Leftrightarrow			
$\exists$	36	\exists	$\forall$	34	\forall	$\neg$	216	\neg
$\vee$	218	\vee	$\wedge$	217	\wedge			
$\spadesuit$	170	\spadesuit	$\heartsuit$	169	\heartsuit	$\diamondsuit$	168	\diamondsuit
$\clubsuit$	167	\clubsuit						

FIG. 4.3 – Liste des symboles spéciaux et leurs identificateurs associés

Symboles spéciaux

On pourra étendre la syntaxe de manière à pouvoir écrire à la place d'un mot un identificateur, désignant un symbole spécial dont la liste est donnée figure 4.3. Cette table donne pour chaque caractère spécial le caractère qui lui correspond dans la fonte Symbol de Postscript, ainsi que l'identificateur qui devra lui être associé.

On pourra ainsi écrire :

```
f(x)\rightarrow +\infty quand x\rightarrow a \Leftrightarrow \forall A \geq 0, \exists \alpha > 0, |x-a| < \alpha \Rightarrow f(x) > A
```

pour produire

$$f(x) \rightarrow +\infty \text{ quand } x \rightarrow a \Leftrightarrow \forall A \geq 0, \exists \alpha > 0, |x - a| < \alpha \Rightarrow f(x) > A$$

Paramétrage du style

Pour pouvoir changer le style d'un document, il est plus pratique de faire précéder le texte d'une commande additionnelle :

```
\documentstyle{nom d'un fichier de style}
```

où un fichier de style contiendra des déclarations de la forme suivante :

```
\textwidth{80}
```

```
\normalfontfamily{Times}
\normalfontsize{3.5}
\titlefontfamily{Helvetica}
\titlefontbold{true}
\titlefontsize{6}
```

Ces fichiers de style doivent pouvoir être lus avec le même analyseur lexical que les fichiers source LaTeX, mais avec un analyseur syntaxique différent.

4.1.6 Interface PASCAL du module POSTSCR

```
(**** fontes *****)

(* les familles de fontes postscript standard. *)
type
  font_family = (Times,Helvetica,Courier,Symbol) ;

(* description d'une fonte. size est la hauteur en mm. *)
(* toute utilisation d'une size <= 0 déclenche une erreur *)
type
  fontdesc = record
    family   : font_family;
    oblique  : boolean;
    bold     : boolean;
    size     : real
  end;

(**** fichiers *****)

(* open_postscript(e,s) ouvre un fichier postscript en ecriture, nomme *)
(* s.ps. Aucune operation d'ecriture n'est possible tant qu'aucun *)
(* fichier n'est ouvert. *)
(* declenche une erreur si l'ouverture ne reussit pas. *)
(* e donne quelle est le codage utilise pour les *)
(* lettres accentuees *)

type
  char_encoding = (ISO_8859, DOS, WINDOWS);

procedure open_postscript(e : char_encoding;
  filename : string);

(* termine le fichier postscript en cours *)

procedure close_postscript;

(**** navigation *****)

(* new_page() enregistre la page courante dans le fichier *)
(* postscript, et demarre une nouvelle page blanche *)

procedure new_page;

(**** text *****)

(* output_text(f,x,y,s) affiche le texte donne par la chaine s sur la *)
(* page courante, a la position (x,y) qui sera le coin inferieur *)
```

```

(* gauche du texte, ceci dans la police de caracteres f. *)
(* x et y sont en millimetres, l'origine etant situee approximativement *)
(* a 2cm du bord gauche et du bas de la feuille *)
(* Declenche une erreur si la position (x,y) déborde de la *)
(* feuille (mais ne verifie pas si le texte lui-meme ne déborde pas !) *)

procedure output_text(f : fontdesc ;
                     x,y : real;
                     s : string);

(* text_width(f,s) retourne la largeur de la chaine s dans la fonte f, *)
(* en millimetres *)

function text_width(f : fontdesc ;
                   s : string):real;

```

À titre d'exemple, voici la suite des instructions PASCAL à exécuter pour créer un fichier Postscript d'une page, contenant un petit texte :

```

uses postscr;

var
  f : fontdesc;
  w : real;

begin
  open_postscript(DOS, 'c:\result');

  f.family := Times;
  f.oblique := false;
  f.bold := false;
  f.size := 10.0 ;

  output_text(f,0.0,230.0,'Voici un texte qui se pour');
  w := text_width(f,'Voici un texte qui se pour');
  output_text(f,w,220.0,'suit ici. ');

  close_postscript;
end;

```

4.2 Les formules mathématiques

Comme première extension, nous souhaitons pouvoir saisir des formules mathématiques qui seront automatiquement composées par le programme.

4.2.1 Syntaxe des formules

Il y aura en fait deux types de formules : les formules simples, dont la hauteur est la même que le texte, qui pourront faire partie du texte d'un paragraphe normal, et les formules complexes, qui formeront un paragraphe à part entière, centré. Les formules simples sont :

- les constantes entières, réelles, par exemple « 23 », « 3.1416 » ;
- les variables, qui seront formées par une lettre majuscule ou minuscule, par exemple « x », « A » ; remarquer que celles-ci seront composées en italique ;

+	75	+	-	77	-	(vide)	*
/	79	/					
±	177	\pm	×	180	\times	÷	184 \div
>	94	>	<	92	<	=	93 =
≥	179	>=	≤	163	<=	≠	185 <>
∈	206	\in	⊂	204	\subset	⊆	205 \subseteq
∪	200	\cup	∩	199	\cap		

FIG. 4.4 – Liste des opérations et leurs identificateurs associés

- les lettres grecques, par exemple « π », « Γ », que l'on saisira `\pi`, `\Gamma` dans le texte source ;
- une opération entre deux formules, par exemple « $x + 1$ » (écrit `x+1` dans le source) , « $E \cup F$ » (écrit `E\cup F`) ;
- une fonction appliquée à une formule ou plusieurs formules, par exemple « $f(x)$ » qui sera écrit simplement `f(x)`, ou bien « $\sin(x)$ » qui sera écrit `\sin(x)` ; noter que la fonction sin est écrite en caractères droits et pas en italique ;
- une formule entre parenthèses « $(x + 1)$ », entre crochets « $[x + y]$ », entre barres « $|a + b|$ » ;
- une formule puissance une autre formule, par exemple « x^2 », ce qui sera noté `x^2` dans le texte source ;
- une formule indice une autre formule, par exemple « x_i », ce qui sera noté `x_i` dans le texte source ;
- une formule puissance une autre formule et indice une troisième, par exemple « x_i^2 », ce qui sera noté `x_i^2` dans le texte source ;

La liste des opérations est donnée figure 4.4, chacune d'elles étant accompagnée du numéro du caractère correspondant dans la fonte Symbol et de sa syntaxe dans le source. Noter que dans la mesure où les parenthèses correspondent à un affichage effectif, les accolades pourront être utilisées pour « parenthéser » les formules, on écrira ainsi `2^{x+y}` pour produire 2^{x+y} et `x_{i+j}` pour produire x_{i+j} , ou encore `x_{i^2}` pour obtenir x_{i^2} .

Les formules complexes seront formées à partir des formules simples par

- des fractions, par exemple

$$\frac{\pi}{2}$$

qui sera noté `\frac{\pi}{2}` dans le source ;

- des signes d'intégration, par exemple

$$\int_0^{+\infty} f(x) dx$$

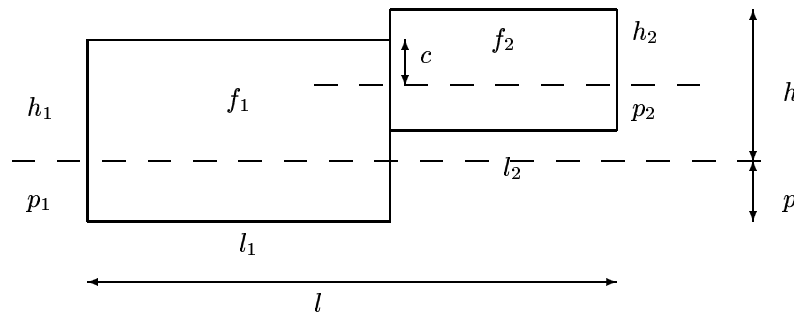
qui sera noté `\int{0}{+\infty}{f(x)*\d x}` dans le source (remarque : si `x` est une variable, alors `\d x` sera une formule valide, composée `dx`) ;

- des signes de sommation, par exemple

$$\sum_{k=1}^n \frac{1}{k}$$

qui sera noté `\sum{k=1}{n}{\frac{1}{k}}` dans le source.

calcul de la taille d'une formule de la forme $f_1 \wedge f_2$ sera calculée à partir des tailles de f_1 et f_2 d'après le schéma suivant :



La distance c entre le haut de la formule f_1 et la ligne de base de f_2 sera calculée par un certain coefficient k_1 multiplié par la hauteur h_2 de f_2 . k_1 sera un paramètre de style, qui pourra être de 0.9 par défaut. La taille de fonte dans laquelle la formule f_2 doit être composée doit être plus petite que la taille de fonte de f_1 , on choisira comme taille de fonte pour f_2 un certain coefficient k_2 multiplié par la taille de fonte de f_1 . k_2 sera un paramètre de style qui vaudra 0.7 par défaut.

Un principe analogue doit être appliqué pour les formules en indice : la ligne de base de la formule en indice de hauteur h_2 doit être de $k_3 \times h_2$ en dessous de la ligne de base de f_1 . k_3 pourra être choisi à 0.5 par défaut. Pour les autres formules complexes, vous devez déterminer des principes analogues.

Enfin, pour les formules simples que sont les variables et les constantes, vous prendrez comme hauteur la hauteur de la fonte, et 0 comme profondeur. la largeur sera calculée à l'aide de la fonction `text_width`.

Ce jugement de taille permet ensuite de définir des règles pour un *jugement de composition* d'une formule, de la forme $x, y, t \vdash f \Rightarrow i$ où dans le contexte, t donne la taille de la fonte et x et y sont les coordonnées d'origine où la formule f doit être composée. i sera une liste d'instructions de composition, qui sont essentiellement des appels à la fonction `output_text`, et occasionnellement à `draw_line`.

4.3 Les tableaux

4.3.1 Syntaxe

Un tableau à m lignes et n colonnes sera décrit dans le texte source sous la forme suivante :

```
\tabular{lcr...}
{case 1,1 & case 1,2 & ... & case 1,n}
...
{case m,1 & case 1,2 & ... & case m,n}
```

Il constitue un paragraphe à part entière, il se termine donc par une ligne vide. Le premier argument de `\tabular` est une séquence d'exactly n lettres parmi `l`, `c` et `r` indiquant que la colonne correspondante sera composée respectivement à gauche, centrée et à droite. le caractère `&` sert de séparateur entre les cases. Une erreur à la composition devra être déclenchée si le nombre de cases données dans l'une des listes de cases n'est pas exactement n . Le nombre de lignes est déterminé uniquement par le nombre de lignes dans le source, et ne peut donc pas donner lieu à une erreur.

Par exemple, la description de haut niveau suivante:

```
\tabular{lcr}
```

```
{Année & semestre & Nombre d'étudiants }
{96-97 & 1 & 27}
{96-97 & 2 & 110}
{97-98 & 1 & 50}
```

devra donner :

Année	semestre	Nombre d'étudiants
96-97	1	27
96-97	2	110
97-98	1	50

4.3.2 Indications pour la composition

Le programme de composition devra composer un tableau en deux étapes : d'abord le calcul de la largeur de chaque colonne et de la hauteur totale du tableau, puis la composition proprement dite. Remarquer que le tableau lui-même devra être centré sur la page. La distance entre le texte des cases et les lignes tracées est un paramètre de style supplémentaire, par défaut 3mm. L'épaisseur des traits est également un nouveau paramètre de style, par défaut 1mm.

Vous devez donc définir un jugement de taille pour les tableaux, et les règles sémantiques associées, ainsi qu'un jugement de composition.

4.4 Calculs automatiques

Pour réaliser la fonction d'un « tableur », nous désirons permettre à l'utilisateur de saisir des formules à calculer au moment de la composition.

4.4.1 Syntaxe

Par convention, ces formules seront tapées entre deux #. Voici un exemple de source souhaité

```
% moy(y) donne la moyenne des 2e et 3e case de la ligne y
\def{#moy(y):=(C(2,y)+C(3,y))/2#}
% sumcol(y) donne la somme des cases c(x,2)+..+c(x,y) ou x est le
% numero de colonne courante
\def{#sumcol(y):=if y<2 then 0 else C(X,y)+sumcol(y-1)#}
% moycol(y) donne la moyenne des cases c(x,2),...,c(x,y-1)
\def{#moycol(y):=sumcol(y-1)/(y-2)#}

\table{1rrr}
{Élève & devoir 1 & devoir 2 & moyenne }
{Agnan & #19# & #19# & #moy(Y)# }
{Alceste & #1# & #2# & #moy(Y)# }
{Clotaire & #0.5# & #0# & #moy(Y)# }
{Eudes & #5# & #8# & #moy(Y)# }
{Rufus & #9# & #10# & #moy(Y)# }
{moyennes & #moycol(Y)# & #moycol(Y)# & #moy(Y)#}
```

et la sortie voulue :

Élève	devoir 1	devoir 2	moyenne
Agnan	19	19	19
Alceste	1	2	1.5
Clotaire	0.5	0	0.25
Eudes	5	8	6.5
Rufus	9	10	9.5
moyennes	6.9	7.8	7.35

Comme pour les formules à composer, les formules à calculer seront analysées avec un analyseur lexical particulier. Le nombre de décimales affichées pour les réels sera un paramètre de style, 2 par défaut.

4.4.2 Sémantique

La sémantique de ces expressions entre $\#$ est informellement définie de la façon suivante :

- dans une instruction `\def`, on définit soit une constante, soit une fonction à argument, par une expression arithmético-logique, éventuellement récursive pour les fonctions ;
- dans un tableau, toute expression entre $\#$ est remplacée par sa valeur, qui est calculée dans un environnement où sont définies :
 1. les fonctions précédemment définies par des `\def` ;
 2. les constantes X et Y qui sont respectivement les numéros de colonne et de ligne de la case en cours de calcul (numérotées en commençant à 1) ;
 3. la fonction spéciale C à deux arguments donnant pour chaque (x, y) la valeur de la case $C(x, y)$ du tableau en cours de composition.

Les expressions autorisées seront les expressions arithmético-logiques avec le `if` et le `let` pour les variables locales. Le programme de composition devra bien entendu vérifier que ces expressions sont bien typées.

Pour éviter des boucles dans le calcul des cases, on considérera que pour calculer la valeur d'une case, on ne peut utiliser que les valeurs des cases d'ordonnée inférieure, ou bien de même ordonnée et d'abscisse inférieure. Les valeurs de chaque case seront donc calculées de haut en bas et de gauche à droite.

D'autre part, on s'interdit de redéfinir une constante ou une fonction, et pour éviter des conflits avec X , Y ou C , on n'autorise comme identificateur de constante ou de fonction uniquement des séquences non vides de lettres minuscules.

Dans votre deuxième rapport, la sémantique informelle décrite ci-dessus devra être formalisée par des jugements de typage et d'évaluation, et des règles sémantiques associées.

Chapitre 5

Un navigateur hypertexte

Dans cette étude de cas, nous allons écrire un programme de navigation dans un *hypertexte*. Un hypertexte est un ensemble de pages de textes (autrement dit des fichiers) qui sont reliées entre elles par des *ancres*, une ancre étant une marque sur un mot du texte qui permet, par sélection de ce mot avec la souris, de faire afficher une autre page de texte.

Un exemple d'hypertexte est donné par les fichiers d'aide de Windows. Un autre exemple à la mode est donné par les fameuses pages du *World Wide Web*, dont la particularité est que les liens peuvent référer à des pages se trouvant physiquement sur une machine située n'importe où dans le monde, à la condition que celle-ci soit reliée au réseau mondial *Internet*.

Sur la toile (mot français pour le *World Wide Web*), les fichiers d'hypertexte doivent respecter un standard international pour que ces pages puissent être visualisées sur n'importe quel type d'ordinateur (PC, Macintosh, station de travail sous UNIX, etc.). Ce standard s'appelle le langage HTML (*HyperText Markup Language*), et dans cette étude nous allons réaliser un outil de visualisation de pages décrites dans ce langage, et de navigation d'une page à l'autre.

Bien sûr, il n'est pas de question d'implanter le langage HTML complet qui comporte trop de fonctionnalités différentes, nous allons implanter les fonctionnalités de base, qui nous permettront de voir les principes et méthodes à utiliser pour implanter toutes les autres fonctionnalités. Nous commençons dans cette première partie avec une version très basique qui sera à étendre dans la deuxième et la troisième partie.

5.1 Le langage HTML de base

Ce langage permet de donner une description de *haut niveau* d'une page de texte : il décrit la structure du texte, c'est-à-dire pour la première partie de notre étude, comment le texte est découpé en paragraphes, mais il ne dit absolument pas où doivent être placés les sauts de lignes, ce sera le travail du programme de visualisation de déterminer ces sauts de ligne.

Un exemple de fichier HTML est donné figure 5.7.

Le programme de visualisation devra afficher le texte de ce fichier en ouvrant une fenêtre graphique et afficher une page comme montré figure 5.8. La chose importante à remarquer est que les espaces entre les mots et les passages à la ligne du texte source n'ont aucune influence sur le texte de sortie : en effet les passages à la ligne sont calculées de manière à ce que les lignes soient contre les marges gauches, en remplissant les lignes au maximum.

Le texte de l'exemple contient également une ancre, qui est définie par la ligne

```
<a href="fichier.htm">ancres</a>
```

```
<p>
Ceci est un paragraphe de texte écrit dans le langage HTML. il commence
par un << p >> écrit entre les caractères << inférieur >> et << supérieur
>>. il se termine par << /p >> écrit entre ces mêmes caractères.
</p>

<p>
Remarquer que les passages à la ligne du texte
initial n'influent pas du tout sur les passages à la ligne du texte
affiché, ainsi que le nombre d'espaces entre les
mots qui sont toujours remplacés par un seul espace
dans le texte affiché.
</p>

<p>
Un fichier HTML peut également contenir des
<a href="fichier.htm">ancres</a> pour pouvoir
naviguer entre les pages de texte.
</p>
```

FIG. 5.1 – Un exemple d'un fichier HTML.

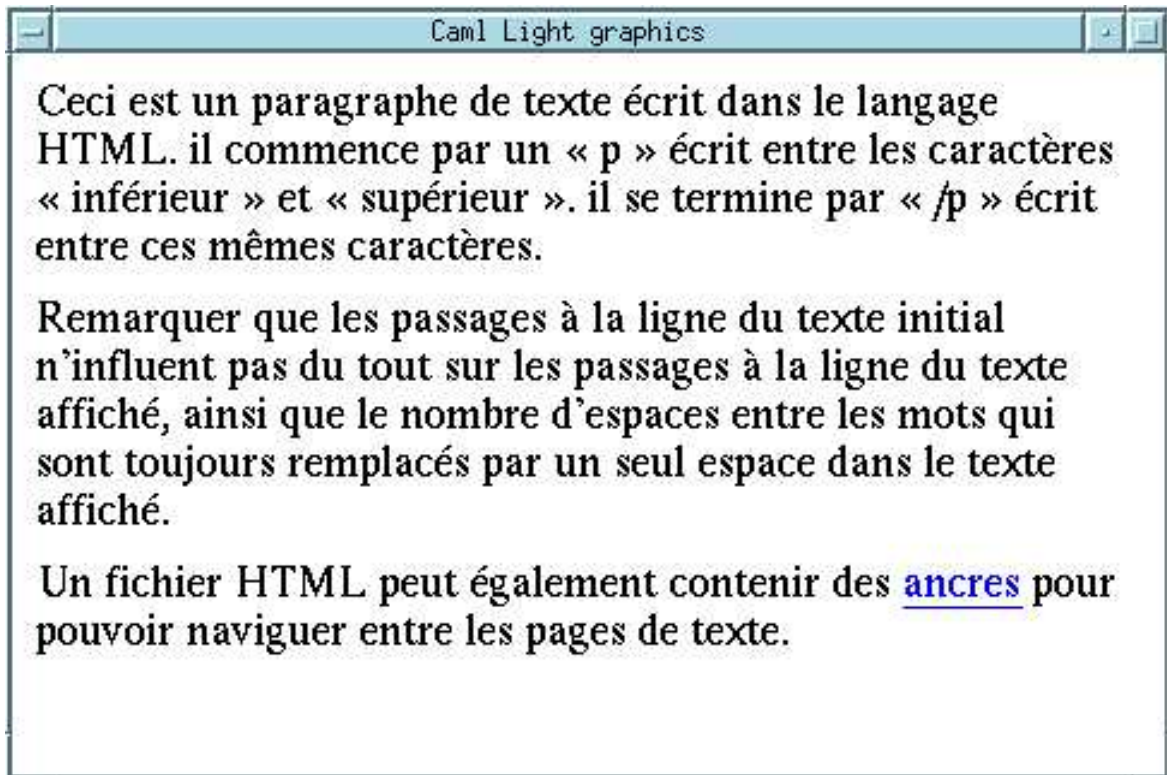


FIG. 5.2 – Le fichier HTML comme affiché par le navigateur.

Dans le texte affiché, le mot sur lequel l'ancre est placé, c.-à-d. « ancrés » dans l'exemple, sera affiché d'une manière spéciale de manière à ce que le lecteur voit que ce mot est marqué par une ancre : il sera souligné et, sur un écran couleur, écrit dans une autre couleur, par exemple bleue. Le *lien* ("fichier.htm" dans l'exemple) n'est pas visible sur la page affichée, mais quand l'utilisateur placera la souris sur le texte de l'ancre et cliquera sur le bouton gauche de la souris, le contenu de la fenêtre graphique sera effacé et la page indiquée par le lien de l'ancre sera affichée.

Le programme se terminera quand l'utilisateur pressera la touche q. Un exemple de navigateur est à votre disposition pour mieux comprendre les fonctionnalités voulues du programme :

- pour le DEUG MIAS : dans `f:\caml\m3mias\html`
- pour la MIAGE : dans `~treinen/exemplaire/partie1`

Description précise du langage

Nous donnons maintenant une description précise du langage. Pour votre premier rapport, vous devrez traduire cette description en une grammaire formelle.

Un fichier HTML consiste en une suite (éventuellement vide) de *paragraphes*. Chaque paragraphe commence par `<p>` et finit par `</p>`. Un paragraphe consiste en une suite (éventuellement vide) d'*éléments*, séparés par un ou plusieurs caractères d'*espacement* (espace, tabulation ou saut de ligne). Un élément d'un paragraphe peut être :

- soit un *mot*, formé d'une suite non vide de caractères autres que les caractères d'espacements, `<`, `>` et `"` ;
- soit une *ancre*, qui est de la forme `<a href=fichier>mot</a>`, où *fichier* est un nom de fichier entre guillemets.

On appellera *fichier* le *lien* et *mot* le *texte* de cette ancre. Le nom de fichier pourra être formé de n'importe quel caractères sauf `"`. On autorisera aussi plusieurs caractères d'espacement entre `<a`, `href`, `=` et le nom du fichier. Les « mots-clés » `a`, `p` et `href` devront être acceptés aussi bien en minuscules qu'en majuscules.

Pour la taille de la fenêtre graphique, des marges et l'espacement vertical entre les paragraphes on choisira des valeurs raisonnables. Quand un fichier est trop long pour être affiché dans la fenêtre graphique, l'affichage s'arrêtera quand la fenêtre est pleine.

Il y a, par contre, des erreurs plus importantes : un fichier peut ne pas exister, ou bien ce fichier peut ne pas être conforme à la syntaxe HTML. Dans ce cas, le navigateur superposera sur la dernière page affichée un petit rectangle d'une couleur de fond différente du texte normal, et contenant le message d'erreur. Le rectangle avec le message d'erreur sera alors supprimé quand l'utilisateur pressera une touche du clavier.

5.2 Énumérations et tableaux

Cette deuxième partie de l'étude est une suite de la première partie. Vous utiliserez pour cette partie les outils CAMLLEX et CAMLYACC présentés en cours, qui nous permettent d'aborder un langage HTML plus riche qu'il n'était possible pour une analyse lexicale et syntaxique programmée avec des automates standard. Pour cette partie, le programme devra obligatoirement être écrits avec des vrais modules compilés.

Le langage HTML considéré dans la première partie de l'étude ne contenait qu'un seul type de paragraphe : une liste de mots et d'ancres, délimités par `<p>` et `</p>`. Le langage HTML étendu de cette deuxième partie de l'étude comprend maintenant deux nouveaux types de paragraphes : les tableaux et les énumérations.

5.2.1 Les tableaux

Un exemple de fichier contenant un tableau est donné figure 5.3, et l’affichage attendu est donné figure 5.4.

De manière générale, un tableau commence par une déclaration de la forme `<table cols="w">` et se termine par `</table>`, où le mot w est la spécification des alignements de chaque colonne. Cette spécification ne doit contenir que les lettres `l`, `c` et `r`, indiquant respectivement une justification à gauche, centrée ou à droite de la colonne correspondante.

Le contenu du tableau est donné ligne par ligne, une ligne commençant par `<tr>`, et dans chaque ligne, chaque case commence par `<td>`. Le contenu des cases est une liste quelconque de mots ou d’ancres.

Une condition sémantique est que toutes les lignes doivent avoir le même nombre de cases, et ce nombre doit être égal à la longueur du mot qui spécifie les alignements des colonnes.

5.2.2 Les énumérations

Une énumération est une liste numérotée de portions de texte, où la numérotation dans le bon ordre doit être réalisée automatiquement par le navigateur.

Un exemple de fichier contenant une énumération est donné figure 5.5, et l’affichage attendu est donné figure 5.6.

Les énumérations peuvent donc être imbriquées et elles peuvent contenir également des tableaux. Les différents points d’une énumération sont numérotés par défaut à partir de 1. Il est aussi possible de spécifier le premier numéro n utilisé par une énumération, dans ce cas-là l’énumération commence par `<ol start=n>`.

D’une manière générale, une énumération commence donc avec `<ol>` ou bien `<ol start=n>` et se termine par `</ol>`. Chaque point d’une énumération commence par `<li>`. Le contenu de chaque point est une suite quelconque de paragraphes simples, tableaux ou énumérations.

5.2.3 Remarques générales sur le langage

Pour simplifier, on n’autorise maintenant les mots clefs (`href`, `start`, etc.) qu’en minuscule. Rappelons que les énumérations et les tableaux peuvent, comme les paragraphes de textes, contenir des ancres qui sont à gérer comme les autres.

Un exemple de fichier HTML montrant tous les extensions considérées dans cette deuxième partie de l’étude est donné figure 5.7, et le résultat attendu est donné figure 5.8.

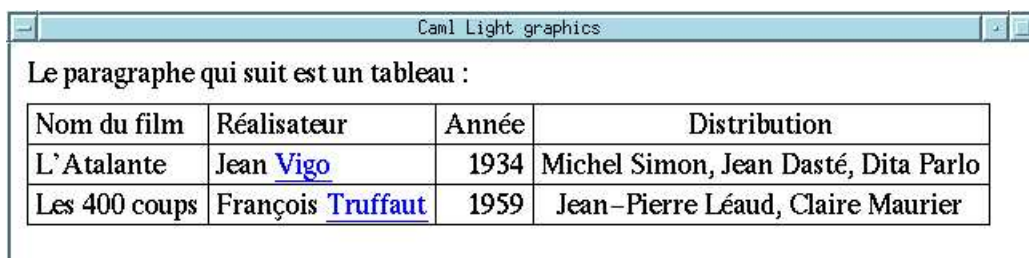
```

<p>Le paragraphe qui suit est un tableau :</p>

<table cols="11rc">
<tr>
  <td>Nom du film
  <td>Réalisateur
  <td>Année
  <td>Distribution
<tr>
  <td>L'Atalante
  <td>Jean <a href="vigo.htm">Vigo</a>
  <td>1934
  <td>Michel Simon, Jean Dasté, Dita Parlo
<tr>
  <td>Les 400 coups
  <td>François <a href="truffaut.htm">Truffaut</a>
  <td>1959
  <td>Jean-Pierre Léaud, Claire Maurier
</table>

```

FIG. 5.3 – Un exemple de tableau.



Le paragraphe qui suit est un tableau :			
Nom du film	Réalisateur	Année	Distribution
L'Atalante	Jean Vigo	1934	Michel Simon, Jean Dasté, Dita Parlo
Les 400 coups	François Truffaut	1959	Jean-Pierre Léaud, Claire Maurier

FIG. 5.4 – Le tableau comme affiché par le navigateur.

```
<p>Le paragraphe qui suit est une énumération</p>

<ol>

<li><p>Ceci est le premier point de l'énumération.</p>

<li><p>Voilà le deuxième point de l'énumération. Noter que la marge
gauche de tout le texte de chaque point de l'énumération est décalé
sur la droite.</p>

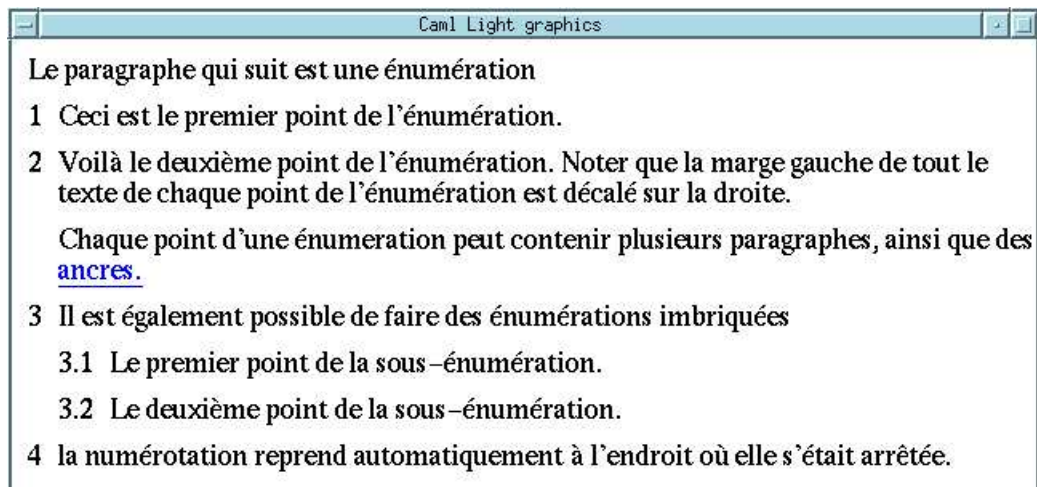
<p>Chaque point d'une énumération peut contenir plusieurs paragraphes,
ainsi que des <a href="ancre.htm">ancres.</a></p>

<li><p>Il est également possible de faire des énumérations imbriquées
</p>

<ol>
<li><p>Le premier point de la sous-énumération.</p>
<li><p>Le deuxième point de la sous-énumération.</p>
</ol>

<li><p>la numérotation reprend automatiquement à l'endroit où elle
s'était arrêtée.</p>

</ol>
```

FIG. 5.5 – *Un exemple d'énumération.*FIG. 5.6 – *L'énumération comme affichée par le navigateur.*

```

<p>Quelques données sur les cours d'Informatique du semestre M3 du
DEUG MIAS</p>

<ol>

<li><p>Le nouveau programme a commencé en septembre 1996</p>

<li><p>Le tableau suivant indique l'évolution du nombre d'étudiants :</p>

<table cols="lcr">
<tr>
<td>Année <td>Semestre <td>Nombre d'étudiants
<tr>
<td>96-97      <td>1          <td>27
<tr>
<td>96-97      <td>2          <td>110
<tr>
<td>97-98      <td>1          <td>50
<tr>
<td>97-98      <td>2          <td>70
</table>
</ol>

<p>On peut faire les remarques suivantes :</p>

<ol start=1997>
<li><p>beaucoup plus d'étudiants au second semestre qu'au premier</p>
<li><p>beaucoup plus équilibré</p>
</ol>

```

FIG. 5.7 – Un exemple d'un fichier HTML avec toutes les extensions.

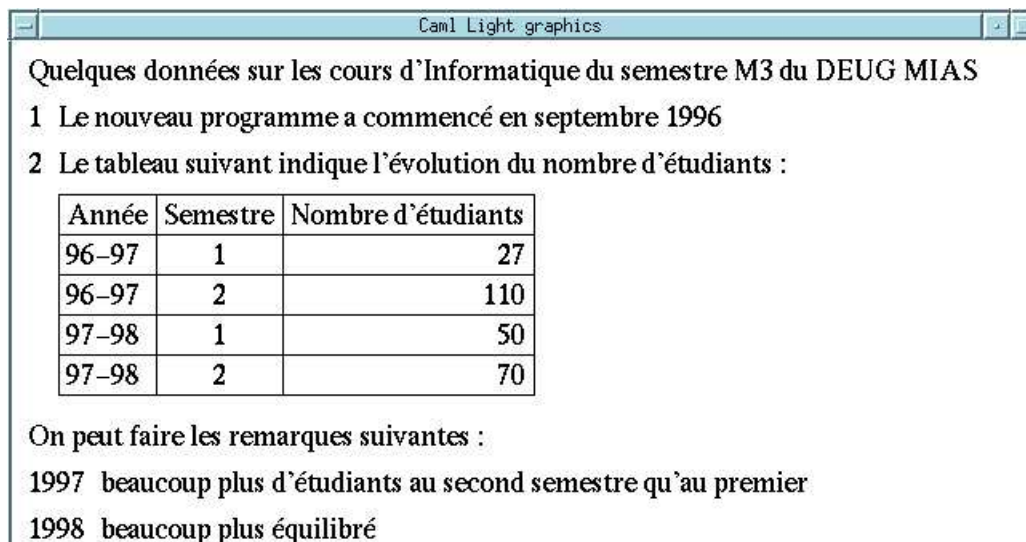


FIG. 5.8 – Le fichier HTML complet comme affiché par le navigateur.

=====

Table des figures

1	Les différentes phases de l'interprétation d'un langage	2
1.1	Deux grammaires décrivant les entiers écrits en binaire	10
2.1	Automate déterministe reconnaissant les entiers naturels en numération binaire.	28
2.2	Automate déterministe complet reconnaissant les entiers naturels en numération binaire.	28
2.3	Automate et grammaires correspondantes pour les entiers naturels en notation binaire.	32
2.4	Un automate non-déterministe.	34
2.5	Automate déterminisé.	35
2.6	Réduction d'un automate non-minimal.	36
2.7	Un automate non minimal.	38
2.8	Automate minimisé.	38
3.1	Un automate non-déterministe avec transitions vides.	50
3.2	Automate non-déterministe sans transitions vides.	51
3.3	Automate déterminisé.	52
3.4	Automate déterminisé et minimisé.	52
3.5	Complémentation de l'automate de la figure 3.4	54
3.6	Automate d'analyse lexicale	58
3.7	Automate déterminisé d'analyse lexicale	58
4.1	Deux arbres de dérivation du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire ambiguë des expressions arithmétiques	74
4.2	Parcours gauche et droit de l'arbre de dérivation du mot $Id + Cte \times Cte + (Id + Id)$ dans la grammaire des expressions arithmétiques désambiguïsée . .	74
4.3	Méthodes d'analyse syntaxique	76
4.4	Analyse syntaxique descendante	76
4.5	Analyse syntaxique ascendante	78
4.6	Shift et Reduce	78
5.1	Comparaison des syntaxes abstraites et concrètes	95
2.1	Règles d'évaluation des expressions arithmétiques	109
2.2	Règles d'évaluation des expressions logiques.	110
2.3	Règles d'évaluation des expressions arithmético-logiques.	112
2.4	Typage des expressions arithmético-logiques	114
2.5	Typage des expressions arithmétiques surchargées	118
2.6	Échec du jugement de typage de $(1 + 3) \times 3.14$	118
3.1	Typage des expressions arithmétiques étendues	122

3.2	Typage de l'expression arithmétique étendue <code>letn = 2 in letm = n + 1 in m + n</code>	122
3.3	Évaluation des expressions arithmétiques étendues	123
3.4	Évaluation de l'expression arithmétique étendue <code>letn = 2 in letm = n + 1 in m + n</code>	124
3.5	Évaluation de l'expression arithmétique étendue <code>letn = 2 in (letn = n + 2 in n) + n</code>	124
3.6	Typage des expressions arithmético-logiques étendues.	126
3.7	Évaluation des expressions arithmético-logiques étendues	127
4.1	Typage des expressions arithmétiques fonctionnelles	132
4.2	Typage de l'expression arithmétique fonctionnelle <code>let funsqr(x) = x × x in sqr(3) + sqr(4)</code> où $\Gamma = \{sqr : \text{Int} \rightarrow \text{Int}\}$	133
4.3	Évaluation des expressions arithmétiques fonctionnelles	133
4.4	Évaluation de l'expression arithmétique fonctionnelle <code>let funsqr(x) = x × x in sqr(3) + sqr(4)</code> où $\Gamma = \{sqr = (x, \times(\text{Var}(x), \text{Var}(x)))\}$	134
4.5	Typage des fonctions récursives	134
4.6	Évaluation de l'expression fonctionnelle récursive <code>let fun recf(x : int) : int = if x = 0 then 1 else x × f(x - 1) in f(3)</code> , partie 1	136
4.7	Évaluation de l'expression fonctionnelle récursive <code>let fun recf(x : int) : int = if x = 0 then 1 else x × f(x - 1) in f(3)</code> , partie 2	137
5.1	Typage des expressions fonctionnelles	141
5.2	Évaluation des expressions fonctionnelles	141
5.3	Codages des expressions <code>let</code>	142
7.1	Typage des programmes impératifs.	150
7.2	Règles d'évaluations des programmes impératifs.	151
7.3	Traduction des programmes impératifs, avec une seule variable x du type <code>int</code> , aux programmes fonctionnels.	152
4.1	Exemple de fichier source simple	165
4.2	Une page de sortie possible pour l'exemple	166
4.3	Liste des symboles spéciaux et leurs identificateurs associés	168
4.4	Liste des opérations et leurs identificateurs associés	171
5.1	Un exemple d'un fichier HTML.	177
5.2	Le fichier HTML comme affiché par le navigateur.	177
5.3	Un exemple de tableau.	180
5.4	Le tableau comme affiché par le navigateur.	180
5.5	Un exemple d'énumération.	181
5.6	L'énumération comme affichée par le navigateur.	181
5.7	Un exemple d'un fichier HTML avec toutes les extensions.	182
5.8	Le fichier HTML complet comme affiché par le navigateur.	182

Index

- accessible, 27
- alphabet, 10
- ambiguë, 48, 51
- arbre de dérivation d'un jugement, 68
- automate
 - avec transitions vides, 34
 - complet, 20
 - déterministe, 20
 - non-déterministe, 26
- axiome (d'une grammaire), 13
- clôture par ε , 35
- compatible, 29
- complémentaire, 11
- complémentation, 39
- concaténation
 - de langages, 11
 - de mots, 10
- conclusion, 66
- dérivation
 - d'un jugement, 67
 - sémantique, 67
- désambiguïsation, 48, 55
- déterminisation, 27
- $\text{Det}(A)$, 27
- domaine
 - de contextes, 65
 - de valeurs, 65
- ε , 10
- $E(A)$, 36
- ε -clôture, 35
- engendré
 - langage, 14
- environnement
 - de typage, 83
 - de valeurs, 85
- état, 20, 26, 34
 - acceptant, 20, 26, 34
 - atteint, 20
 - initial, 20, 26, 34
 - poubelle, 21
- étiquette
 - abstraite, 58
- étoile de Kleene, 11
- exécution, 20
 - d'un automate non-déterministe, 26
 - d'un automate avec transitions vides, 35
- expression
 - abstraite, 58
 - de syntaxe abstraite, 58
- fonction
 - de transition, 20
- grammaire, 12
 - abstraite, 58
 - hors contexte, 13
 - régulière, 18
 - régulière réduite, 19
- hors contexte, 13
- intersection, 11
- itéré, 11
 - strict, 11
- jugement, 65
 - d'évaluation, 85
 - de typage, 83
- Kleene, 11
- langage, 11
 - dénoté par une expression rationnelle, 40
 - engendré par une grammaire, 14
 - fortement typé, 74, 78
 - régulier, 18
 - rationnel, 40
 - reconnaissable, 20
- lettre, 10
- longueur d'un mot, 10
- mot, 10
- non-terminal, 13

palindrome, 16
portée, 85
poubelle, 21
prémisses, 66
produit de concaténation, voir concaténation
puissance
 d'un langage, 11
 d'un mot, 10

règle
 de grammaire, 13
 sémantique, 66
reconnu
 par un automate, 20
réécrire, 14

source, 13
substitution, 100
symbole, 10
 d'entrée, 13
 non terminal, 13
 terminal, 12

table des symboles, 42
terminal, 12
token, 41
transition, 26, 34
 fonction de, 20
 vide, 34
typage
 fort, 78
type, 83

union, 11

vocabulaire, 10, 20, 26, 34
 non-terminal, 13
 terminal, 12