

Du pourquoi au comment : la théorie de la démonstration de 1950 à nos jours

Jean-Yves Girard

Institut de Mathématiques de Luminy, UPR 9016 – CNRS
163, Avenue de Luminy, Case 930, F-13288 Marseille Cedex 09

girard@iml.univ-mrs.fr

« *A la fin du siècle dernier, les mathématiques étaient gravement menacées par les paradoxes ; mais la théorie de la démonstration put redonner un sens aux mots de la tribu.* »

Cette citation imaginaire résume l'idéologie moyenne du théoricien de la démonstration de 1950. Elle situe d'emblée la théorie de la démonstration dans une problématique fondamentaliste (l'élimination des paradoxes) qui affirme que la logique donne le sens profond des mathématiques, ce que j'appellerai le « pourquoi ». Plus tard, vers 1985, l'informatique devait promouvoir une approche plus pragmatique, ce que j'appellerai le « comment » : ce *comment* est une préoccupation bien moins noble que le *pourquoi*, mais qui demande un appareillage beaucoup plus subtil ¹.

Le domaine s'organise autour de la scission entre une aile conservatrice et idéologique en récession (quoique toujours influente) et une aile novatrice et (parfois trop) pragmatique... cette scission n'est pas nouvelle et on peut en trouver les linéaments dans la controverse sur les fondements qui opposa Hilbert et Brouwer ² dans les années 1920, une controverse qui se termina par l'expulsion de Brouwer de la rédaction des *Mathematische Annalen*, voir [38].

1. La France continentale est connexe *parce qu'on* peut lier n'importe quelle ville à Paris ; mais ne pas mépriser la question « *comment* la France est-elle connexe ? » requiert de construire un réseau de communication beaucoup moins trivial qu'une simple étoile centrée sur Paris.

2. Bien que Brouwer ne soit pas strictement un logicien et encore moins un théoricien de la démonstration, son *intuitionnisme* nous apparaît aujourd'hui comme l'ancêtre du *comment*.

1 Préhisteoire du pourquoi

1.1 Problèmes de fondements

La théorie des ensembles vaut avant tout par sa capacité d'unification : elle énonce fondamentalement l'unité des mathématiques ³. L'outil de cette unité, c'est la construction progressive des rationnels, des réels, etc. à partir des entiers, qui peuvent à leur tour être considérés comme des cardinaux. Ce qui ne veut pas dire qu'un nombre réel soit vraiment un ensemble de rationnels, et encore moins un sous-ensemble de $\mathbb{N}$: la vraie nouveauté c'est la possibilité de combiner librement tous les aspects du raisonnement mathématique... après tout, d'où tenons-nous que les méthodes analytiques et les méthodes algébriques de la théorie des nombres ne conduisent pas à des résultats divergents ?

A la fin du siècle dernier la théorie des ensembles *naïve* se ramène pour l'essentiel au *schéma de compréhension*, i.e. le fait que toute propriété A définit un ensemble

$$\exists X \forall x (x \in X \leftrightarrow A[x]) \quad (1)$$

Quand en 1897 Burali-Forti [4, 5, 39], publie son paradoxe ⁴, il n'en résulte qu'une bien vague menace —non sur les mathématiques, mais sur les ponts entre diverses parties des mathématiques—. Et d'ailleurs, il ne fallut que 10 ans (Zermelo 1908, voir [39]) pour remarquer qu'une restriction —somme toute empirique— du schéma de compréhension aux éléments d'un ensemble Y préexistant

$$\exists X \forall x (x \in X \leftrightarrow A[x] \wedge x \in Y) \quad (2)$$

plus quelques cas autres cas particuliers de (1) (ensemble des parties, ensemble des entiers) ainsi que l'*Axiome du Choix* permet de réaliser l'unité des mathématiques sans contradiction apparente. 90 ans de mathématiques formalisables à l'intérieur ⁵ de la théorie de Zermelo ont confirmé la parfaite viabilité de ces restrictions.

Alors pourquoi cette grande peur de l'an 1900 ? Hilbert avait-il vraiment peur quand il se dressa en chevalier blanc pour sauver l'édifice des mathématiques ou ne faisiat-il qu'instrumentaliser un prétendu danger pour promouvoir une extrême forme de scientisme, le *formalisme* ?

En 1900, —quinze ans avant l'ypérite— la science n'avait encore montré que son aspect Jules Verne et la foi *scientiste* ne se connaissait aucune borne. D'où l'idée de fonder les mathématiques une fois pour toutes ; de par la nature du *credo* scientiste, cette fondation ne pouvait être que mathématique, et pour éviter l'auto-justification, elle devait donc prendre la forme d'une *réduction* des mathématiques

3. A opposer à la physique, formée d'ilôts reliés par des passerelles hasardeuses.

4. Appelons *ordinal* un ensemble bien ordonné par la relation d'appartenance ; l'*ensemble* des ordinaux étant lui-même un ordinal, il doit s'appartenir, et donc n'est pas bien ordonné.

5. Sauf la théorie des ensembles qui échappe de par sa thématique à cette limitation.

abstraites (théorie des ensembles) à un corpus de mathématiques très élémentaires (arithmétique à l'ancienne).

Ce qu'on a appelé le *Programme de Hilbert* remonte à la fameuse liste de problèmes de 1900, et après une première tentative en 1904 [23, 39], prend corps dans les années 20, voir [24, 39]. On peut en donner deux versions :

- Les démonstrations de cohérence : démontrer qu'un système formel ne mène pas à contradiction.
- Les démonstrations de conservation : démontrer qu'un système formel ne démontre pas plus de résultats élémentaires que les méthodes « à la Papa ».

De plus, pour établir l'un ou l'autre de ces buts, seules les méthodes élémentaires de démonstration sont permises.

Il est facile de voir au moyen de raisonnements élémentaires que les deux versions sont en fait équivalentes : ainsi la cohérence n'est que la conservation par rapport à une impossibilité élémentaire genre $0 = 1$. L'aspect « conservation » prolonge d'ailleurs toute une tradition de théorie des nombres : donner des démonstrations élémentaires, e.g. éliminer l'analyse complexe dans le théorème des nombres premiers. Sous cette forme Hilbert énonce donc une conjecture de « pureté des méthodes ».

1.2 Le refus du réalisme

Les fondements se doivent de répondre à l'angoissante question : quand je dis que A est vrai, qu'est-ce que ça veut dire ? La *vérité* est de fait une notion rebelle à toute analyse, car elle commute à toutes les opérations logiques ($A \wedge B$ est vrai ssi A et B le sont, $\neg A$ est vrai A n'est pas vrai etc.) et en particulier toute justification des principes mathématiques par leur vérité est grandement suspecte. Pour s'en convaincre, considérons l'arithmétique de Peano, qui ne contient guère —outre des principes logiques visiblement vrais— que le schéma de démonstration par récurrence, dit *schéma d'induction* ; or si $A[0]$ est vrai et si $A[n] \Rightarrow A[n+1]$ est vrai pour tout n , on se convainc aisément que $A[n]$ est vrai pour tout n ... mais on vient précisément de justifier l'induction sur A au moyen d'une induction sur la vérité de A : on tourne en rond. On est ainsi amené à se méfier du réalisme et, par exemple, pour montrer que l'équation en f, g

$$f(g(x)) = f(f(x)) \tag{3}$$

ne conduit pas à $f(x) = g(x)$, on remarquera —plutôt que d'exhiber des fonctions f, g *ad hoc*— que les conséquences de (3) comportent f simultanément aux deux membres.

On a été ainsi amené à éliminer toute référence à la vérité, en essayant de la remplacer par la *prouvabilité*. Et il est de fait que la prouvabilité de A dans un système formel fixé (arithmétique, théorie des ensembles) a une autre allure que sa vérité. Le programme de Hilbert aurait d'ailleurs justifié complètement cette substitution, au moins pour des énoncés de structure simple, Π_1^0 (voir 1.3). Il est quand même nécessaire de remarquer que l'élimination de la vérité dans les textes de théorie de la démonstration atteint souvent des sommets... d'hypocrisie. Ainsi, quand on étudie l'équation (3), on peut très bien construire un domaine à deux éléments 0, 1 et les fonctions $f(0) = f(1) = 1$, $g(0) = 0$, $g(1) = 1$ et le calcul de vérité *fini* qui montre que (3) est vrai sans que $f(0) = g(0)$ le soit exactement la même valeur épistémologique qu'une récurrence sur les enchaînements d'équations ⁶. Il ne faut jamais nommer le diable et en théorie de la démonstration traditionnelle, le diable c'est le monde extérieur.

1.3 Le théorème de Gödel

Au départ Gödel a sans doute fait sienne la confusion entre vérité et prouvabilité, et essayé d'obtenir une contradiction dans les mathématiques : comme on sait —dans un système $\mathcal{T}$ contenant un minimum d'arithmétique— coder et définir rigoureusement la démontrabilité, on peut par un argument de diagonalisation ⁷ produire un énoncé G qui exprime formellement sa non-prouvabilité, autrement dit sa fausseté, ce qui mène à contradiction. Minute, papillon... pourvu que la prouvabilité soit identique à la vérité, ce qui fait de G un énoncé du type « *cré-tois* », i.e. « *Je mens* ».

Il reste la possibilité que la vérité soit distincte de la prouvabilité. A ce moment-là il faut être soigneux. On classe les énoncés arithmétiques en Σ_n^0 et Π_n^0 : $n \geq 1$ majore le nombre d'alternances de quantificateurs de A dont le premier est donné par le choix du symbole $\Sigma = \exists$ ou $\Pi = \forall$. Ainsi le théorème de Fermat, la conjecture de Riemann, les énoncés de cohérence sont Π_1^0 , alors que $L(1, \chi) \neq 0$, la prouvabilité d'un énoncé sont Σ_1^0 , tandis que $\forall \chi L(1, \chi) \neq 0$ est Π_2^0 . Les énoncés Σ_1^0 vrais sont toujours démontrables (pour démontrer $\exists n A(n)$ avec A sans quantificateurs, il suffit de trouver n et d'effectuer le calcul de vérité pour $A(n)$, qui étant fini, se formalise), et par dualité les énoncés Π_1^0 démontrables dans $\mathcal{T}$ sont vrais dès que $\mathcal{T}$ est cohérente. L'énoncé de Gödel est Π_1^0 , et s'il était démontrable, sa démontrabilité (i.e. sa négation) serait démontrable en temps que Σ_1^0 vrai, c'est à dire que G et $\neg G$ seraient démontrables, ce qui rendrait $\mathcal{T}$ contradictoire. On en déduit donc que G n'est pas démontrable, ce qui veut dire que G est vrai. C'est le premier *théorème d'incomplétude* (1931) qui réfute la forme « conservation »

6. Ce comportement se retrouve même chez Gentzen, [12, 14] : un banal argument de vérité dans un modèle fini est transcrit au moyen d'expressions cabalistiques et acquiert à ce moment le statut de démonstration de cohérence.

7. Utilisé par Cantor pour démontrer que $\mathcal{P}(\mathbb{N})$ n'est pas dénombrable et déjà démarqué par Russell en 1901 dans son fameux paradoxe, voir [39], pp. 124-125.

du programme de Hilbert : en effet prenons pour $\mathcal{T}$ un système d'arithmétique contenant exactement les méthodes « élémentaires », et qui est donc forcément cohérent. . . l'énoncé G résulte de la cohérence de T sans pouvoir être obtenu dans $\mathcal{T}$, i.e. par des méthodes élémentaires.

En fait, un peu plus de travail (mais beaucoup plus de soin), montre que G peut être remplacé par la cohérence de $\mathcal{T}$: c'est le deuxième théorème d'incomplétude, obtenu dans la foulée. Il réfute la version « cohérence » du programme de Hilbert. Ce théorème est devenu une tarte à la crème, délicieusement employée à contre-emploi⁸. Il détruit (ou du moins aurait dû détruire) tout espoir de réaliser le fameux programme de Hilbert : pour fonder $\mathcal{T}$ il faudra utiliser plus que $\mathcal{T}$, ce qui ne convaincra jamais que les vrais croyants.

1.4 Le refus de l'évidence

Le premier théorème d'incomplétude interdit toute forme de commutation de la prouvabilité à la négation et induit dans les systèmes logiques un « trou » entre les théorèmes (énoncés démontrables) et les antithéorèmes (énoncés réfutables), trou qu'on ne saurait « boucher » de façon raisonnable⁹. On retrouve cette situation dans de nombreux problèmes d'algorithmique : ainsi tout algorithme prétendant résoudre le 10^e problème de Hilbert¹⁰ sera soit incomplet, soit fautif (Matiasevič 1970). Plus abstraitement, le problème d'arrêt d'un programme n'est pas décidable, c'est à dire qu'étant donné un programme $P[.]$ dépendant d'un paramètre n il n'y a pas en général de programme $Q[n]$ qui puisse nous dire si oui ou non $P[n]$ s'arrête.

Sautons hardiment 50 ans : vers 1980, certains spécialistes de l'intelligence artificielle s'émeuvent qu'on ne sache pas toujours répondre « oui » ou « non » à toute question, et proposent donc de compléter la logique de façon à répondre dans tous les cas, tout ceci en contradiction formelle avec le théorème d'incomplétude et ses corollaires, qu'ils ignorent ou nient. On a ainsi assisté à l'éclosion *ex nihilo* d'une paralogique¹¹ qui se maintient dans la plus parfaite nullité, mais sans disparaître totalement, du fait de la prégnance d'un nouveau type de scientisme lié à l'ordinateur qui serait —pensent-ils— capable de répondre à tous les problèmes.

Ainsi ne faut pas s'étonner que le théorème de Gödel n'ait pas été compris à l'époque : tout le monde, y compris Gödel, a cherché des portes de sortie permettant de contourner la difficulté.

8. Et a même fait l'objet d'un best-seller d'une kolossale vulgarité : *Gödel-Escher-Bach*.

9. De la même façon qu'on ne saurait « compléter » un opérateur non-borné sur un Hilbert $\mathbb{H}$.

10. La résolution des équations diophantiennes

11. « Logique » non-monotone, « logique » des défauts etc.

1.5 Gentzen

Les années 1930 sont dominées par deux grands théoriciens de la démonstration, tous deux morts prématurément : Herbrand (1908-1931) et Gentzen (1909-1945).

1.5.1 Le calcul des séquents

Le théorème d’Herbrand, [21, 22, 39] (1930) anticipe le théorème de Gentzen (souvent appelé *Hauptsatz*¹²) dont il constitue une version « synthétique ». Bien que le théorème d’Herbrand ait certains avantages sur le résultat de Gentzen, il est beaucoup moins souple d’emploi et nous nous concentrerons sur le *Hauptsatz*. Ce résultat de 1934, voir [10, 14], est la plus parfaite réalisation du programme de Hilbert : par exemple c’est un résultat de pureté des méthodes qui montre (c’est la *propriété de la sous-formule*) que pour démontrer un énoncé A , on peut se restreindre aux sous-énoncés de A ; par ailleurs le résultat est établi par des méthodes finitistes, i.e. élémentaires. Où le bât blesse c’est que le résultat ne s’applique qu’à la logique pure (i.e. le *calcul des prédicats*) : il ne fonctionne plus (ou alors seulement avec des béquilles) dès qu’on ajoute ne serait-ce qu’un peu d’arithmétique, i.e. des axiomes d’induction. Tel quel, le résultat ne s’applique au programme de Hilbert que pour prouver la cohérence de systèmes formels ridiculement faibles.

Le théorème de Gentzen, voir annexe B, domine la théorie de la démonstration qui est d’ailleurs loin d’en avoir épuisé toutes les significations. Du point de vue technique les règles logiques sont réécrites au moyen d’un *calcul des séquents*, basé sur de profondes symétries. Une règle (dite de coupure) permet de donner à ce calcul la souplesse déductive : elle exprime de façon adroite la transitivité de la conséquence logique, c’est à dire la possibilité d’utiliser dans une démonstration des résultats intermédiaires (lemmes). Gentzen produit un algorithme (élimination des coupures) qui permet de remplacer effectivement une démonstration par une démonstration sans coupures, c’est à dire pratiquement directe, non-déductive. De telles démonstrations n’existent pas dans la nature (bien que de nos jours l’ordinateur puisse en reconstituer) pour des raisons de taille et surtout de compréhension : comme il est impossible de redescendre du général au particulier on démontrera dix fois la même propriété dans des cas similaires plutôt que de démontrer la propriété-souche dont ces dix cas découlent. Ces démonstrations ont des propriétés absolument remarquables qui compensent leur caractère artificiel. On peut rapprocher le calcul des séquents de la mécanique hamiltonienne : même insistance sur les symétries, et même inadaptation aux problèmes concrets compensée par une hauteur de vue remarquable.

1.5.2 Démonstration de cohérence

En 1936 Gentzen s’attaque à l’arithmétique de Peano, un système qui contient beaucoup plus que toutes les méthodes « élémentaires », et qui est vraiment soumis au théorème de Gödel. Sur la base de ses travaux sur le calcul des séquents, il

12. Ce mot allemand ne veut rien dire de plus que l’expression *théorème fondamental* utilisée par Herbrand pour son propre résultat.

en donne en 1938 [13] une seconde démonstration de cohérence... par induction transfinie jusqu'à l'ordinal dénombrable ϵ_0 ¹³. Il s'agit *grosso modo* de théorèmes d'élimination des coupures imparfaits, mais suffisants pour assurer la cohérence. On reste confondu par la disparité entre l'ingéniosité du travail et la nullité du résultat, du moins du résultat « officiel »: « *Gentzen a établi la cohérence de l'induction jusqu'à ω au moyen de l'induction jusqu'à ϵ_0* »¹⁴.

Le première démonstration de cohérence de Gentzen (1936) [11, 14], est encore plus injustifiable du point de vue des fondements: Gentzen y développe ni plus ni moins une interprétation interactive des démonstrations, vues en tant que stratégies pour établir la vérité des formules. Ce travail qui violait un des interdits majeurs (la notion de stratégie gagnante contient le prédicat de vérité) n'a pas eu à l'époque de postérité... Ce n'est qu'avec l'ère du « comment » et l'abandon du fondamentalisme qu'on a trouvé la perspective correcte: il s'agit de la première interprétation *ludique* de la logique. Incidemment (nous n'aurons pas le temps d'y revenir) la théorie des jeux est devenue, à la fin du siècle, en association avec la logique linéaire, un des axes majeurs de la théorie de la démonstration.

2 Le pourquoi

La poursuite des travaux de Gentzen dans une optique aveugle de démonstrations de cohérence domine le sujet jusqu'en —disons— 1985. J'appellerai ceci la théorie du « *pourquoi* », voir note 1. A ce propos le théorème de Gödel avait rendu les protagonistes plus ou moins schizophrènes: par exemple quand en 1972 je déclare naïvement à Schütte (reprenant une opinion de Kreisel, voir plus bas) qu'une démonstration de cohérence de la théorie des ensembles n'aurait pas la moindre valeur épistémologique, il me répond: « *D'accord mais je me sentirais quand même mieux si j'en voyais une* ». Il est d'ailleurs à noter que malgré le théorème d'incomplétude et le caractère fastidieux des « démonstrations de cohérence », il y a toujours un public pour cela, qui cherche sans doute à être rassuré; on voit ainsi des informaticiens —par ailleurs raisonnables— s'inquiéter de démonstrations de cohérence et réagir comme Schütte quand on leur dit qu'elles ne prouvent rien.

Bien entendu, il faut faire la différence entre cohérence de la théorie des ensembles qui se heurte au théorème de Gödel et la question de la cohérence de l'axiome du choix (résolue par Gödel en 1938). Ce résultat (comme ceux de Cohen en 1963) ne se heurte en aucune façon au théorème d'incomplétude, vu que la cohérence de la théorie des ensembles sans axiome du choix est supposée. Il s'agit donc d'une réduction élémentaire de la théorie avec axiome du choix à la théorie sans cet axiome (la réduction étant faite au moyen d'une traduction explicite de l'une

13. Le premier point fixe de la fonction $x \rightsquigarrow \omega^x$.

14. Méchanceté d'un grand mathématicien français rapportée par Kreisel. Nous aurons l'occasion de revenir sur ce point.

dans l'autre). Autrement dit le théorème de Gödel s'oppose aux démonstrations de cohérence absolue, mais pas à des résultats de cohérence *relative*.

2.1 Schütte et l'école de Munich

Vers 1950 la théorie de la démonstration se concentre pour l'essentiel en Allemagne, mentionnons Ackermann, Lorenzen et celui qui fera (tardivement) école, Schütte. On y démontre toujours des cohérences de l'arithmétique au moyen du sempiternel ϵ_0 , mais les méthodes ont un peu changé. Schütte se permet des règles infinies du genre

$$\ll \text{De } A[n] \text{ pour } n = 0, 1, 2, \dots \text{ déduire } \forall x A[x] \gg$$

pour laquelle il démontre un excellent analogue du *Hauptsatz* [33], les ordinaux comme ϵ_0 intervenant pour mesurer la hauteur des arbres de démonstration.

C'est avant tout un progrès sur l'hypocrisie qui caractérisait certains articles de Gentzen. Nous connaissons l'insistance de Gentzen à rester fini (par exemple ses ordinaux sont remplacés par leurs formes normales de Cantor, expressions finies), en particulier à éviter toute notion de démonstration infinie. Il est pourtant clair que Gentzen a pensé ses travaux en fonction des démonstrations infinies, mais qu'il a préféré recouvrir ses traces pour ne pas prêter le flanc à l'accusation d'abandonner le fini. Or il est facile de voir que de telles « démonstrations » peuvent être vues comme des arbres à embranchements dénombrables, dont on peut, en commençant par la conclusion, décrire effectivement des portions finies arbitraires : ce qui fait que les dogmes finitaristes ne sont jamais violés.

L'approche de Schütte permet donc une notable simplification de la théorie de la démonstration, en particulier l'extension de la propriété de la sous-formule à l'arithmétique, mais sur une base de logique infinitaire. Cette machinerie a été appliquée à des extensions de l'arithmétique de Peano au second ordre, où l'on admet la quantification sur des ensembles d'entiers : de telles formules sont classées en Σ_n^1 et Π_n^1 , où ici $n \geq 1$ compte l'alternance des quantificateurs ensemblistes. Il y a une foultitude de tels systèmes¹⁵, qui se distinguent par la complexité logique des formules A sur lesquelles le schéma de compréhension (2) est admis (ainsi on est toujours resté très en dessous de deux quantificateurs ensemblistes, i.e. Π_2^1) et accessoirement par des finasseries sur le principe d'induction transfinie. Des travaux américains des années 60 (Feferman, Friedman, Howard, Tait) ont permis d'écrire ces divers systèmes sous forme de *définitions inductives* : par là on veut dire que les ensembles en question sont obtenus par certaines itérations dénombrables. La première extension notable est due à Takeuti, [35], (1967) qui traite de la compréhension à un quantificateur ensembliste (i.e. Π_1^1) au moyen d'un système d'ordinaux de son cru (*diagrammes ordinaux*), assez peu compréhensible.

15. D'autant plus que les progrès sont lents et qu'il faut bien publier.

C'est dans la thèse de Jane Bridge (une synthèse de différents travaux, notamment Bachmann et Aczel, 1972), voir [2], que l'école de Munich va puiser son échelle de mesure ordinale. Le travail de Bridge (rebaptisé « système de Buchholz ») permet aux élèves de Schütte : Pohlers, Buchholz, Jäger d'analyser diverses définitions inductives en termes de divers ordinaux. C'est une technique bien rodée qui marche sans surprise, ni mauvaise (c'est correct) ni bonne (c'est toujours la même sauce), voir [3].

Cette activité continue toujours, quoique de façon beaucoup plus modeste. Il faut bien entendu rappeler que l'objection fondamentale qui nous vient de Gödel n'est en aucune façon contournée par ces travaux : ainsi, à théorie plus grosse, ordinal plus gros etc. On a donc vu la création de catalogues mettant en relation des théories formelles assez artificielles avec des ordinaux assez gigantesques... il y a de la *Kabbale*, de la numérologie là-dedans, car ces résultats n'apportent en général aucune information intéressante. Ce qui n'est pas tout à fait vrai pour les résultats de Gentzen et ses extensions immédiates, qui se font en termes d'ordinaux « raisonnables » : ainsi Friedman, voir e.g. [9], a-t-il remarqué que l'ordinal associé à une petite extension de Peano est en bien un ¹⁶ du fait d'un célèbre théorème dû à Kruskal sur les plongements d'arbres finis, et donc le théorème d'incomplétude interdit toute démonstration du théorème de Kruskal dans ce système.

Bien entendu la notion de démonstration infinie pose problème : il faut éviter le gag qui consisterait à toujours remettre le début de la démonstration, c'est à dire l'existence de branches infinies. C'est précisément pour éviter de telles branches que les démonstrations sont plongées dans des ordinaux comme ϵ_0 . La cohérence des systèmes résulte alors de la bonne fondation des ordinaux utilisés. La bonne fondation, la notion de démontrabilité en logique infinie s'expriment par des formules Π_1^1 : cette classe de formules joue pour la logique infinie le rôle joué par la classe Σ_1^0 pour la logique finie. Mais ce n'est pas tout : même si on sait l'approximer effectivement par des portions arbitrairement grandes, même si on l'a plongée effectivement dans un ordinal, comment s'assure-t-on de tout ça ? Il nous faut bien une démonstration finie auxiliaire de tout ça... ça donne lieu à beaucoup de codage sans beaucoup d'imagination, pour à la fin nous laisser sur une impression d'inachevé.

2.2 Kreisel

Kreisel domine sans conteste la théorie de la démonstration de cette époque. Sa période d'activité proprement dite s'arrête vers 1970, mais son influence court encore pendant une bonne dizaine d'années. Son attitude fondamentale est de refuser ou de minimiser la valeur des démonstrations de cohérence (point de vue

16. Cet ordinal Γ_0 est présenté comme un ordre linéaire, et la question est de savoir s'il s'agit d'un bon ordre ; du bon ordre de Γ_0 résulte la cohérence des théories dont les démonstrations infinies peuvent être plongées dans Γ_0 .

qui apparaît rétrospectivement comme le seul tenable) pour n'en retenir que la valeur ajoutée, autrement dit « *Qu'est-ce que ça nous apprend ?* » : il y a déjà un début de passage au *comment* chez Kreisel. Donnons quelques exemples de sa méthodologie :

- La démonstration de cohérence la plus stupide « *Les axiomes de l'arithmétique sont vrais, les règles préservent la vérité* » est répudiée avec horreur par tous les spécialistes, car elle justifie les axiomes par eux-même (on y croit) et utilise un prédicat de vérité par nature tautologique et infini. Cela ne rebute pas Kreisel qui n'en considère que l'aspect formel : il se restreint à un nombre fini d'axiomes d'induction et la propriété de la sous-formule permet alors de *formaliser* cet argument dans l'arithmétique elle-même, qui prouve ainsi la cohérence de ses sous-systèmes finis, et qui donc par le théorème d'incomplétude ne saurait être finiment axiomatisée. C'est ce qu'on appelle le *schéma de réflexion*, voir [28] ou [17], ch. 4.
- Une démonstration de cohérence se fait toujours au moyen d'un algorithme d'élimination des coupures (voir section 4) qui énumère donc les fonctions entrée/sortie du système associées à des énoncés $\Pi_2^0 \forall n \exists m A[n, m]$. Autrement dit la démonstration permet d'analyser les fonctions récursives (algorithmes) dont le système prouve la terminaison : c'est le thème des fonctions *prouvablement totales* : la démonstration de cohérence permet de caractériser ces fonctions, en particulier leur croissance.
- De la même façon une démonstration de cohérence ordinale permet de borner les ordinaux récursifs dont le système prouve la bonne fondation : c'est le thème des *ordinaux prouvables*. Encore une fois une information peut être extraite de la démonstration.

La critique de Kreisel a été plus récupérée qu'assimilée ; ainsi les articles de théorie de la démonstration à la Papa ne se terminent-ils plus sur « *La théorie $\mathcal{T}$ est cohérente* », mais sur « *Les ordres récursifs prouvablement bien fondés de $\mathcal{T}$ sont ceux plus petits que α* » où α est l'ordinal utilisé dans la démonstration de cohérence. Il ne faudrait d'ailleurs pas voir dans Kreisel un opposant systématique à la théorie de la démonstration germanique, car au fond il ne rêvait que d'assoir le domaine sur une base moins idéologique. Il a ainsi repris l'idée de méthodes finitistes étendues (c'est déjà chez Gentzen) en nuancant le mot cruel rapporté plus haut (note 14) : l'induction jusqu'à ω dans l'arithmétique de Peano est permise sur des énoncés de complexité logique arbitraire, alors que l'on peut restreindre l'induction jusqu'à ϵ_0 utilisée dans la démonstration de cohérence à des énoncés très élémentaires (sans quantificateurs).

L'exigence de Kreisel était avant tout la demande de plus de substance mathématique. Il avait en fait commencé très tôt à *appliquer* la théorie de la démonstration à l'analyse effective de démonstrations de théorie des nombres, e.g. au théorème de

Littlewood concernant les changements de signe de $\pi(x) - \text{Li}(x)$. La possibilité *a priori* de ce travail vient de la complexité logique du théorème de Littlewood (Π_2^0) : une des contributions théoriques essentielles de Kreisel, [25, 26], est la remarque que les théorèmes Π_2^0 classiques sont aussi démontrables de façon intuitionniste, i.e. ont un contenu effectif ¹⁷, voir 3.1 Il a essayé d'entraîner la théorie de la démonstration dans cette direction d'applications à la théorie des nombres (analyse logique de théorèmes non-effectifs), mais sans véritablement convaincre ; il faut bien dire que ce type d'application qui est basé sur des remplacements explicites inspiré du théorème de Gentzen devient facilement fastidieux et incompréhensible, par saturation de données explicites que les mathématiques usuelles ont le bon goût de cacher. Par exemple, si on s'intéresse à la théorie analytique des nombres on ne peut guère couper au problème d'expliciter le prolongement analytique d'une fonction au moyen d'un chemin recouvert de petits cercles etc.

2.3 Les dilatateurs

C'est une expérience isolée (1975-1985), inspirée des critiques de Kreisel. En cherchant à mettre un peu de mathématiques dans ce fatras de codages ordinaux en tous genres, j'ai remarqué que non seulement les ordinaux sont des limites directes d'entiers (ce qui est évident), mais que la plupart des constructions commutent aux limites directes, ce qui leur donne un côté finitaire immédiat, et aussi aux produits fibrés, ce qui donne des résultats de forme normale. La théorie des *dilatateurs* (foncteurs des ordinaux dans eux-même préservant limites directes et produits fibrés) a été développée essentiellement par moi-même pendant 10 ans, voir [15] en insistant sur les aspects géométriques absents dans l'école de Schütte (e.g. démonstration infinie obtenue par limite directe de démonstrations finies etc.)[16]. Les outils mis en œuvre pouvaient sans doute traiter le cas Π_2^1 , inaccessible aux méthodes à la Schütte ¹⁸, mais à quoi bon ? Au niveau où la différence se manifeste, les mathématiques sont très raréfiées et on ne trouve plus guère que de la théorie des ensembles assez abstraite (avec laquelle d'ailleurs les dilatateurs s'entendent bien). Il fallait retomber sur du concret, du simple : les considérations géométriques à l'œuvre dans les dilatateurs ont engendré la logique linéaire, un sujet d'intérêt beaucoup plus central : ainsi se termina l'expérience, dont il ne reste guère qu'un beau résultat : le théorème de comparaison de hiérarchies de [15] qui relie la hiérarchie paresseuse de Hardy H_α aux autres hiérarchies de fonctions calculables : l'indice α apparaît comme la limite directe des $H_\alpha(n)$.

17. Cela devient faux pour les énoncés du type « *L'équation (E) n'a qu'un nombre fini de solutions* », qui sont Σ_2^0 .

18. La notion d'une démonstration de largeur ordinale arbitraire mais entièrement déterminée par les cas finis est beaucoup plus puissante que le cas de largeur ω considéré par Schütte : on passe de la complexité Π_1^1 à la complexité Π_2^1 tout en étant beaucoup plus « fini ».

2.4 Et maintenant ?

Ce type de théorie de la démonstration n'a pas disparu ; les épigones de Schütte continuent d'aligner des listes de théories en face de *Panzerdivisionen* ordinales sans percée ni conceptuelle (expliquer ce qui se passe) ni en complexité logique (la compréhension Π_2^1).

Une certaine arrière-garde a trouvé à se recycler dans le préchi-précha prédicativiste. Rappelons qu'à l'origine la prédicativité est une idée de Poincaré (reprise par H. Weyl) pour expliquer les paradoxes : on n'aurait pas le droit de définir un objet en termes d'un ensemble qui le contient¹⁹. Il se trouve que —Poincaré ou pas— c'est une idée d'une remarquable stérilité. Comme produit de substitution aux démonstrations de cohérence, on trouve donc des systèmes prédictifs qui sont proposés comme « plus sûrs » que les systèmes habituels. Ça n'a guère de sens, d'abord car le doute quant à la cohérence doit être soit universel, soit répudié comme douteux²⁰, ensuite car les soi-disant méthodes non-prédicatives sont surtout une facilité. Finalement ces théories prédicativement correctes, c'est un peu comme un régime à la carotte, c'est très ennuyeux, mais au fond ça ne fait ni bien ni mal... on pourra consulter [8], un de ces articles où l'idéologie s'exerce au dépend du contenu technique, cantonné à un ramassis de banalités.

Plus récemment des théoriciens de la démonstration plus orientés vers la complexité algorithmique ont essayé de faire revivre les techniques traditionnelles (surtout autour de l'arithmétique), comme Buss avec son *arithmétique bornée* [6]. C'est une ouverture intéressante et ça aurait même place dans la section sur le *comment* si les méthodes n'étaient pas tant figées. En particulier, on aurait bien besoin d'une théorie de la démonstration de l'arithmétique (bornée ou pas) qui ait des outils un peu tranchants ; bien sûr c'est difficile, car tout semble s'opposer à un vrai théorème d'élimination des coupures (finitaire), mais les portes auxquelles on ne frappe pas ne s'ouvrent pas non plus.

3 Préhistoire du comment

La théorie de la démonstration consiste pratiquement toujours en un algorithme de réécriture dont on démontre la terminaison (algorithme d'élimination des coupures, normalisation etc.). Cette terminaison induit des corollaires comme la cohérence. Ce qui veut dire que la seule question de la terminaison de l'algorithme concentre toutes les limitations liées au théorème de Gödel, et donc qu'un obsédé du *pourquoi* (démontrer la cohérence) ne s'intéressera qu'à la terminaison. Maintenant, faisons notre deuil une fois pour toutes des ambitions réductionnistes, fondamen-

19. Kreisel faisait remarquer malicieusement que le plus petit entier vérifiant une certaine propriété est aussi défini en termes d'un ensemble qui le contient.

20. Kreisel : « Les doutes quant à la cohérence sont plus douteux que la cohérence elle-même ».

talistes etc., et admettons le bien-fondé des mathématiques usuelles, dans lesquelles on sait bien montrer la terminaison de l'algorithme ²¹. Alors notre algorithme se termine... c'est tout? Non, il peut avoir des propriétés, et même des propriétés remarquables. Commencer à étudier les propriétés de l'algorithme de réécriture en finesse et non pas en force, c'est passer au *comment*, voir encore la note 1.

Bien entendu la question du *comment* n'est pas récente en théorie de la démonstration; mais c'est seulement dans les années 80 qu'elle a pris le pas sur celle du *pourquoi*, à cause de l'informatique qui a remis à plat certaines questions. On s'est alors aperçu que l'étude du calcul des séquents de Gentzen avait été négligée: on n'y voyait qu'un module, une boîte noire, nécessaire à la théorie de la démonstration infinie. En regardant de plus près, on y découvre un monde d'une richesse inouïe.

3.1 La tradition constructive

Il y a en théorie des nombres une tradition de théorèmes effectifs; ainsi savoir qu'une propriété vaut pour des entiers suffisamment grands est beaucoup plus faible que d'avoir une estimation sur la taille des exceptions possibles, surtout au temps de l'ordinateur qui peut ratisser des segments initiaux assez conséquents. La différence effectif/non-effectif n'est d'ailleurs pas une différence superficielle, c'est une différence logique: ainsi, il suffit de vérifier une propriété $P[n]$ jusqu'à un entier N bien choisi pour qu'elle soit vraie partout (prendre pour N la première exception à P s'il y en une, 0 sinon)... ce truisme sans intérêt ne peut avoir aucun contenu effectif, vu que j'ai pu l'énoncer sans rien savoir de P . Le refus de telles « non-constructions » est à la base de l'*intuitionnisme* de Brouwer et du système logique qui en découle, qui —rappelons-le— n'admet pas le *tiers-exclu* $A \vee \neg A$.

Brouwer nous a laissé une idée remarquable: une démonstration est —ou plutôt devrait être— la construction d'un objet. Par exemple il explique une démonstration de $A \vee B$ comme la construction d'un objet qui vérifie A ou d'un objet qui vérifie B (d'où la réfutation du tiers-exclu). Ce qu'il ne faut pas interpréter de façon stupide: une démonstration de $A \vee B$ ne saurait être obtenue en démontrant soit l'un soit l'autre: qui donc —sauf un cryptographe— énoncerait $A \vee B$ s'il a démontré A ? L'exigence de Brouwer c'est —qu'en principe du moins— on puisse transformer une démonstration de $A \vee B$ en une démonstration qui en soit une de A ou une de B . De la même façon une méthode sera effective parce qu'elle nous donne à défaut d'une valeur numérique, une façon (algorithme) de l'obtenir.

Malgré l'opposition style Don Camillo/Peppone entre le mystique orientalisant Brouwer et le scientifique Hilbert, ces deux point de vue se sont progressivement rapprochés: en effet les démonstrations intuitionnistes sans coupures (qui sont

21. Ce qui fait problème c'est de la montrer avec des moyens limités.

des créations artificielles et donc la remarque de bon sens qui précède ne s'aurait s'appliquer ²²) sont directement explicites. En résumé la logique intuitionniste produit des objets (e.g. des valeurs numériques) implicites que l'élimination des coupures peut expliciter.

Sous son aspect *constructiviste* (bâtir une contre-mathématique) l'intuitionnisme a fait long feu (dernières étincelles dans les années 60 avec... Kreisel); par contre la *constructivité* —c'est à dire l'utilisation hors de tout esprit sectaire des acquis de l'intuitionnisme— reste d'actualité comme nous allons le voir.

3.2 Le lambda-calcul

C'est dans la constructivité qu'il faut insérer l'invention technique du lambda-calcul; à ce propos, il y a deux types d'erreur scientifique: celles qui ne donnent rien et celles qui ont une postérité inattendue et le λ -calcul est de ces dernières. Vers 1930 Church, Curry, Kleene, Rosser, etc. ont l'idée saugrenue de construire une théorie naïve des fonctions (voir annexe A), où toute fonction peut s'appliquer à toute autre et où toute expression définit une fonction. Si on se rappelle qu'un ensemble peut être vu comme une fonction caractéristique, cela n'aurait dû être qu'une nouvelle mouture de la théorie naïve des ensembles. Et d'ailleurs le paradoxe de Russell (qui est la construction d'un point fixe pour la négation) se transpose immédiatement et mécaniquement de façon à obtenir un point fixe pour toute fonction ²³. Ce qui sauve le λ -calcul du ridicule c'est qu'il y a une porte de sortie, i.e. la possibilité que les fonctions soient partielles. Le lambda-calcul devient alors une théorie extrêmement souple des algorithmes partiels; il s'agit même d'une algorithmique universelle. Malgré tout le lambda-calcul reste un domaine un peu marginal jusque dans les années 60, où il commence à acquérir ses lettres de noblesse, grâce en particulier à Böhm, voir [1, 29].

4 Le comment

4.1 Autour de Prawitz

L'activité mathématique de Dag Prawitz, philosophe de profession, donc peu technicien, ne s'étend que sur quelques années (autour de 1968). Il ressuscite un système (déjà considéré par Gentzen) la *déduction naturelle*. Dans le cadre intuitionniste, ce système fait aussi bien que le calcul des séquents; en fait beaucoup mieux, car si le calcul des séquents a un algorithme d'élimination des coupures assez sale (bourré de choix arbitraires), la déduction naturelle supprime des distinctions in-

22. Une démonstration intuitionniste sans coupure de $A \vee B$ est une démonstration de A ou une démonstration de B . Il est donc possible de faire commuter la prouvabilité et la disjonction —contrairement à la négation.

23. Le $n^{\text{ième}}$ avatar de l'argument diagonal de Cantor.

essentielles, de façon à quotienter naturellement le calcul des séquents. Ainsi, la déduction naturelle a une propriété de Church-Rosser, i.e. donne un calcul déterministe.

En s'appuyant sur des travaux antérieurs de Heyting, Kolmogoroff (~1930) et Curry, Howard propose en 1969 ce qu'on appelle l'*isomorphisme de Curry-Howard*, voir annexe A, entre les démonstrations en déduction naturelle et une variante du lambda-calcul (« simplement typé »). Il s'agit de bien plus qu'une bijection, puisqu'il y a transfert de structures déjà existantes : ainsi l'implication intuitionniste $A \Rightarrow B$ devient l'espace des fonctions de A dans B , la conjonction $A \wedge B$ devient le produit cartésien de A et B , la règle de coupure devient la composition des fonctions et la normalisation de la déduction naturelle devient le calcul formel sur les fonctions.

En 1970 j'introduis le système $\mathbb{F}$, voir annexe A, qui est à la fois une déduction naturelle et un lambda-calcul du second ordre, voir par exemple [18]. On n'est déjà plus dans le *pourquoi* : en effet, la convergence des calculs dans le système $\mathbb{F}$ implique formellement la cohérence de *toute* ²⁴ l'arithmétique du second ordre, mais ce résultat de terminaison n'est obtenu qu'à l'aide de méthodes ensemblistes qui ne se veulent en aucune façon élémentaires.

Dans le même ordre d'idées, Martin-Löf crée vers 1972 sa théorie des types [31] qui reste du premier ordre, tout en introduisant d'autres degrés de liberté.

4.2 Le lambda-calcul typé

Ce n'est qu'aujourd'hui qu'on voit ces travaux dans une optique moderne, non réductionniste. Du point de vue réductionniste la terminaison des calculs dans le système $\mathbb{F}$ était un aboutissement sans postérité. C'est une compilation sur le système $\mathbb{F}$ qui devait en faire sortir l'intérêt informatique au début des années 80. En effet ce système, malgré un nombre restreint (cinq) de primitives permet de définir les « types de données » courants (entiers, listes, booléens, arbres etc.) au moyen de formules logiques. L'implication entre deux types de données $D \Rightarrow E$ est le type des algorithmes totaux envoyant des entrées D sur des sorties E . On voit que les formules logiques jouent le rôle de spécifications (quelle type de donnée est accepté, ce qu'il en advient) et garantissent l'absence de boucle (terminaison). De plus la présence d'objets du second ordre permet de donner une grande souplesse aux spécifications : par exemple si $\text{liste}(\alpha)$ est la donnée correspondant aux listes (suites finies) d'objets de type α , le renversement de listes pourra être exprimé sans référence à α comme un objet de type $\forall \alpha (\text{liste}(\alpha) \Rightarrow \text{liste}(\alpha))$, c'est ce qu'on appelle le *polymorphisme*. Enfin un foncteur d'oubli associe aux expressions

24. Rappelons encore que les méthodes pseudo-effectives à l'aide d'ordinaux sont bloquées à la complexité Π_2^1 .

du système $\mathbb{F}$ un lambda-terme (dans le vieux lambda-calcul de Church) qui a le même comportement algorithmique. Il s'agit du développement d'un concept de programmation *fonctionnelle*, en opposition au paradigme dominant de programmation *impérative* ²⁵.

Dans la lignée du système $\mathbb{F}$ on a vu l'éclosion de divers lambda-calculs typés, comme le *calcul des constructions* de Coquand, voir [7], une synthèse du système $\mathbb{F}$ et de la théorie des types de Martin-Löf résolument orientée vers l'informatique. Les ambitions sont clairement affichées : établir un système de programmation fonctionnelle efficace, et pour l'instant, il n'existe que des produits pré-industriels, comme COQ (développé autour de Gérard Huet à l'INRIA Rocquencourt). *Grosso modo* on propose d'extraire des programmes à partir de démonstrations mathématiques : on démontre l'existence d'une solution à un problème et on va traiter la démonstration (compilation) pour en faire un algorithme de calcul d'une solution au problème. L'avantage de cette méthode c'est qu'elle produit des programmes mathématiquement certifiés. Mais les problèmes à résoudre sont multiples :

- D'abord il faut des démonstrations constructives (intuitionnistes), ce qui force à se placer dans un univers mathématique où certains des théorèmes les plus simples ne sont plus vérifiés, à moins, pour les énoncés Π_2^0 d'utiliser la remarque de Kreisel, mais on se heurte alors à des problèmes ardu de conversion classique/intuitionniste ²⁶.
- Et puis il faut pouvoir écrire formellement une démonstration ; ça suppose la mise sur pied de systèmes d'abréviations formelles en tout genre ²⁷ géré par ordinateur.
- Ensuite il faut extraire un programme ; en principe le lambda-terme associé donne un code certifié.

Tout ça est simpliste ; ainsi le programme extrait de la démonstration de connexité de la France (voir note 1) susmentionnée est le réseau SNCF centré sur Paris que nous connaissons... ce n'est pas très efficace. C'est pourquoi on a tendance à compliquer le schéma ci-dessous : rien ne remplace une véritable idée algorithmique et on a donc proposé (en particulier J.-L. Krivine) de partir d'un algorithme (obtenu au « pifomètre »), puis de démontrer qu'il répond à la question et de continuer comme ci-dessus. Le lambda-terme final est une espèce de compilation certifiée de l'algorithme de départ. On voit que le thème du lambda-calcul typé s'est orienté

25. La programmation impérative utilise des instructions genre « del » qui indiquent les actes à effectuer, ici effacer un registre, et qui ne correspondent en aucune façon à un calcul fonctionnel.

26. Une démonstration classique ne définit un algorithme que *modulo* une réduction non-déterministe à un système intuitionniste ; en cela la logique classique n'est jamais constructive, car le contenu algorithmique des preuves n'est pas implicite dans la démonstration.

27. Ainsi mon théorème de normalisation pour le système $\mathbb{F}$ a-t-il été entièrement formalisé sur machine par Berardi ; les parties conceptuelles n'ont pas donné trop de mal, par contre les parties « évidentes » ont donné beaucoup de fil à retordre.

résolument hors du fondamentalisme, jusqu'à devenir (ou du moins chercher fureusement à devenir) appliqué.

Peut-être que la théorie n'y trouve pas tout à fait son compte... C'est à voir ; par exemple les *opérateurs de mise en mémoire* de Krivine [30] utilisent certaines démonstrations intuitionnistes de façon à modifier le type d'exécution d'un algorithme (e.g. forcer l'évaluation d'un argument avant tout calcul).

Il faut faire un sort au paradigme de la *programmation logique*²⁸. Au départ une idée simple : « *Posez le problème et PROLOG fera le reste* ». C'est encore le résultat de Gentzen, mais exploité dans une voie orthogonale : on peut réduire l'idée de calcul à celle de démonstration (c'est implicite dans la démonstration du théorème de Gödel) dans un système très simple qui vérifie le *Hauptsatz*. On se pose donc le problème de démontrer automatiquement certaines formules, ce qui est plausible à cause de la propriété de la sous-formule qui restreint drastiquement l'espace de recherche, voir annexe B ; d'autre part l'aspect logique garantit contre toute erreur. Il y a une limitation évidente : l'algorithme ne peut pas toujours converger, sinon on pourrait décider la prouvabilité en désaccord avec le théorème de Gödel et ses corollaires. Il y a surtout une évidence technologique : on propose-là un algorithme universel, une espèce de véhicule tout terrain fatalement inférieur à des algorithmes utilisant des idées²⁹. Il aurait été sage de confiner PROLOG à des algorithmes de type « exploration », e.g. compiler des fichiers de police... Las, on a voulu mettre des moustaches à la logique, en rajoutant à la recherche des « instructions de contrôle » permettant au programmeur d'utiliser son astuce. Le résultat c'est qu'on ne contrôle plus rien du point de vue logique, et que le slogan de départ s'avère une fumisterie. Pourtant l'idée méritait mieux : il lui aura manqué un peu de modestie — se cantonner à un certain type de tâches — et puis surtout des théoriciens... il est sidérant de penser que la théorie de la programmation logique tient presque entièrement dans les œuvres de Herbrand et Gentzen, voir e.g. [18].

4.3 La sémantique dénotationnelle

Revenons au lambda-calcul, dont la version typée est isomorphe (par Curry-Howard) à la logique intuitionniste. Bien que très robuste, il se présente comme un calcul fonctionnel formel, sans interprétation « naïve ». Plus précisément, le calcul de Church (non-typé) n'admet aucune interprétation ensembliste pour des raisons de circularité (e.g. la fonction identique vérifie $I(a) = a$ pour tout a , y compris $a = I$), tandis que le calcul typé admet bien une interprétation ensembliste triviale dont le cardinal explose bien au delà du continu, en contradiction avec la finitude locale

28. Qui a eu son heure de gloire avec l'engouement des industriels Japonais dans les années 80, maintenant retombé au profit d'une arnaque radicale, la « logique » floue.

29. Si je pose logiquement un problème de tri, il est impossible que l'application de techniques génériques de démonstration automatique me donne une efficacité du type « quicksort ».

des calculs.

On peut voir le lambda-calcul typé comme une *catégorie cartésienne fermée*, où $\text{Hom}(X, Y)$ est lui-même un objet de la catégorie, ce qui s'exprime au moyen d'un certain nombre d'isomorphismes canoniques. Où donc trouver de telles catégories (hormis la désespérante catégorie des ensembles) ?

L'idée d'une interprétation topologique est naturelle ; pourtant il va falloir munir l'espace $\text{Hom}(X, Y)$ des applications continues de X dans Y d'une topologie, et il n'y a pas de miracle : on se heurte aux exigences contradictoire de la convergence simple et de la convergence uniforme. C'est pourquoi la solution proposée par Scott (et aussi Ershov) en 1969, voir e.g. [34], n'est topologique que de nom : les ouverts des *domaines de Scott* sont les segments finaux de certains treillis. Ces topologies sont seulement $\mathcal{T}_0$, —la forme la plus faible de séparation— et non-uniformisables, ce qui évite d'avoir à choisir entre les deux topologies sur l'espace des fonctions. Quoiqu'il en soit, même si l'aspect topologique est un peu douteux, cette *sémantique dénotationnelle* a le bon goût de rester de cardinal raisonnable (au plus le continu) et de fonctionner (grâce à une approximation par limite directe) dans le cas non-typé, voir [1].

La sémantique dénotationnelle n'est pas un simple graphe associé aux fonctions ; elle prend en compte un aspect algorithmique, i.e. la dépendance entrée/sortie de bribes d'information. Typiquement, elle distinguera entre la fonction constante par accident $f(n) = n - n$ et la fonction constante par vocation $f(n) = 0$ qui ont pourtant le même graphe. Autrement dit la sémantique dénotationnelle permet —de façon limitée mais bien réelle— de distinguer plusieurs algorithmes au-dessus d'une même fonction. Ainsi, il été possible de travailler sur la notion d'algorithme *séquentiel* au moyen de considérations purement dénotationnelles ; la condition de *stabilité* de Berry (1978) est satisfaite par la sémantique dénotationnelle des algorithmes séquentiels.

La modélisation du système $\mathbb{F}$ pose un nouveau problème, du fait des types polymorphes, et là l'approche de Scott trouve sa limite naturelle. Par exemple la fonction « retournement » de type $\forall \alpha (\text{liste}(\alpha) \Rightarrow \text{liste}(\alpha))$ prend comme argument un type σ pour devenir ensuite le retournement des listes de type σ . Pour éviter la circularité, j'ai essayé, en m'inspirant des dilatateurs de dire que le retournement est en fait défini sur les seuls domaines finis, puis étendu par limite directe à des types quelconques. Malheureusement un domaine de Scott n'est pas limite directe de domaines finis... et il faut changer la sémantique, en introduisant les *espaces cohérents* : il s'agit de domaines de Scott tellement simplifiés qu'ils sont induits par des graphes : la condition de stabilité de Berry apparaît alors comme la préservation de certains produits fibrés.

4.4 La logique linéaire

Qui peut le plus peut le moins ; les espaces cohérents sont aussi un modèle de la logique du premier ordre. Ils sont suffisamment simples pour qu'on puisse s'amuser à faire des calculs explicites (ce qui est plus problématique avec la sémantique de Scott, très redondante). Ces calculs font vite apparaître le besoin d'abréviations pour des configurations intermédiaires non-logiques... ou plutôt qui n'ont pas de statut logique. Par exemple, les espaces cohérents développent une algèbre linéaire tout à fait décente, et la notion de fonction linéaire de X dans Y prend un sens tout à fait naturel. En fait, si on interprète les règles de la logique intuitionniste en algèbre linéaire³⁰, e.g. en interprétant l'implication comme l'espace des applications linéaires, on découvre que presque toutes les règles du calcul des séquents passent le test, à l'exception de deux principes :

- l'*affaiblissement*, qui permet d'introduire des hypothèses fictives (si B alors $A \Rightarrow B$) et qui doivent se traduire par des dépendances fonctionnelles fictives $f(x) = b$: la linéarité ne suffit pas à cela, il faut admettre des fonctions affines.
- la *contraction*, qui permet de réutiliser les hypothèses (si $A \wedge A \Rightarrow B$, alors $A \Rightarrow B$), i.e. j'ai le droit d'utiliser A deux fois pour démontrer B et de ne le compter qu'une fois. Ici, il nous faut à partir d'une fonction —au départ bilinéaire— $f(x, x')$, fabriquer une fonction d'un seul argument, qui ne peut être que $f(x, x)$, qui se trouve donc quadratique... on sort ainsi violemment de la linéarité³¹

Or il se trouve que Gentzen a mis ces deux principes dans un groupe à part de règles —qu'il a appelées *structurelles*, pensant sans doute qu'elles allaient de soi³²— qui ne dépendent pas des connecteurs logiques. En particulier, on peut modifier la logique en omettant ces règles : on obtient ainsi le premier groupe de règles de la *logique linéaire*. A cette occasion on remarque que la conjonction se laisse décliner suivant deux modes : *multiplicatif* (produit tensoriel $A \otimes B$) et *additif* (somme directe $A \& B$ ³³). Quant à l'implication elle devient implication linéaire, notée $A \multimap B$. La différence entre les deux conjonctions résulte de la gestion de la linéarité : de $A \multimap B$ et $A \multimap C$ on peut déduire $A \multimap B \& C$, mais seulement $A \otimes A \multimap B \otimes C$ ³⁴. Les répétitions d'hypothèses (exprimées par le tenseur) ne sont pas gratuites, et on voit ainsi que l'implication linéaire correspond à une utilisation unique de l'hypothèse... mieux, $A \multimap B$ exprime qu'à partir de A j'ai B , mais

30. Ce n'est pas de l'algèbre linéaire métaphorique : on sait maintenant interpréter la logique dans les espaces de Banach.

31. Dans les espaces de Banach, il faut utiliser les fonctions analytiques sur la boule unité pour « récupérer la contraction ».

32. Il y en a une troisième, l'*échange*, qui énonce la commutativité de la logique (on peut permuter les hypothèses) et dont le rejet au profit d'une logique non-commutative est problématique.

33. Sur les espaces de Banach, normée ℓ^∞ ; la norme ℓ^1 sur la somme directe définit le dual du connecteur $\&$, i.e. $A \oplus B$.

34. De ce point de vue on peut voir la logique linéaire comme le remise en cause de la commutation de la prouvabilité à la conjonction.

que je ne garde pas A . C'est donc une vision *causale* de la déduction logique, qui s'oppose à la pérennité de la vérité traditionnelle en philosophie et en mathématiques. On a ici des vérités fugaces, contingentes, dominées par l'idée de *ressource* et d'action.

Le deuxième groupe de connecteurs est constitué par le point d'exclamation $!A$ ³⁵ qui représente l'algèbre symétrique; cette construction permet de linéariser des fonctions multilinéaires par changement formel de domaine, et donc d'accepter l'affaiblissement et la contraction pour les formules de la forme $!A$: l'implication usuelle devient donc $!A \multimap B$, et on peut voir $!A$ comme la pérennisation de A , i.e. que les ressources sur A sont potentiellement infinies³⁶. L'isomorphisme $!(A \& B) \simeq !A \otimes !B$ suggère le nom d'*exponentiel* pour ce groupe.

Enfin, *last but not least*, la négation linéaire, basée sur l'idée d'espace dual³⁷, induit une involution qui associe à chaque connecteur un connecteur miroir. Concrètement la négation correspond à la dualité « action/réaction » et pas du tout à l'idée de ne pas effectuer une action: typiquement lire/écrire, envoyer/recevoir, sont justiciables de la négation linéaire.

La logique linéaire est vraiment une théorie de la démonstration du comment, i.e. toute en finesse:

- La logique linéaire colle à la dynamique et permet de représenter le comportement de machines abstraites au moyen de la prouvabilité, ce que la logique habituelle ne sait pas faire sauf à caparaçonner les configurations dans une monstrueuse gangue temporelle. De ce point de vue, on se rapproche de la programmation impérative, ce qui est exploité dans une version linéaire de la programmation logique (déplacer des points sur un écran, modifier leur couleur etc.) qui a donné lieu à un premier développement logiciel (Jean-Marc Andreoli, pour Xerox à Grenoble).
- La négation linéaire nie en fait la directionalité, i.e. la distinction hypothèse/conclusion, ou encore entrée/sortie. Ce qui fait de la logique linéaire un outil fondamental pour l'étude du parallélisme asynchrone. Il y a des mathématiques non-triviales là-dedans, en particulier, la première approche géométrique aux démonstrations, les *réseaux de preuve*.

4.5 Et demain?

Il n'y a guère de doute que les applications à l'informatique vont se développer. On attend quand même toujours de vrais produits industriels, mais il y a quelques

35. Et son dual $?A$.

36. Ce qui est le cas pour la vérité mathématique, qui ne s'use pas.

37. Dans les Banach, il faut se donner le « dual ».

frémissements, comme nous l'avons vu...

Quant aux développements théoriques, ils seront avant tout d'ordre sémantique³⁸. La théorie de la démonstration doit se dépouiller de toute syntaxe, de toute bureaucratie³⁹. Ainsi la sémantique dénotationnelle n'est pas complètement satisfaisante, car non refermée sur elle-même : on aurait besoin de théorèmes de complétude, c'est à dire de justifier les règles logiques sans le moindre recours à un « espace de vérité », quel qu'il soit ; cette question et les questions voisines de séquentialité, de *full abstraction* sont très porteuses. Il en est de même des questions de sémantique dynamique qui peuvent prendre l'aspect de *sémantique des jeux*—dont le précurseur est l'incontournable Gentzen, voir section 1.5.2— de *machines abstraites*, par exemple celle de Krivine, ou celui plus mathématique de *géométrie de l'interaction*, [19] formulée en termes d'algèbre stellaires. L'unification mathématique de toutes les formes de sémantique est le problème principal qui se pose à nous.

Ce qui justifie cette entreprise, c'est la découverte progressive de structures cachées dans ce qui ne se présentait *a priori* que comme du pur *symbol pushing*, ce qui tend à accréditer l'idée que les démonstrations ne sont que la verbalisation de structures physiques préexistantes douées de dynamique, laquelle est par ailleurs codifiée par le théorème de Gentzen. Cette hypothèse nous permet d'espérer que des ponts sérieux, pourraient s'établir avec la physique, et spécialement la physique quantique.

Deux annexes techniques

A Le lambda-calcul

Le lambda-calcul est un calcul formel de termes, obtenus à partir de variables $x, y, z, \dots$ par les opérations d'application $t(u)$ et d'abstraction $\lambda x t$. Dans l'application tout terme peut être indifféremment fonction t ou argument u (de même qu'en théorie des ensembles tout objet est indifféremment ensemble ou élément) ; dans l'abstraction, la variable x est muette, et il faut comprendre cette expression comme l'application qui à a associe $t[a/x]$, i.e. le résultat du remplacement de x par a dans t (de même qu'en théorie naïve des ensembles, on peut former $\{x; P\}$ l'ensemble des x vérifiant P). Le lambda-calcul est donc un démarquage fonctionnel de la théorie naïve des ensembles, et le schéma de compréhension (1) trouve son strict analogue dans l'équation

$$\lambda x t(u) = t[u/x] \quad (4)$$

38. Une analyse sémantique du mot « sémantique » ferait apparaître la presque absence de signification de ce mot ; on ne l'utilise ici que par commodité, par opposition à la vieille tradition syntaxique.

39. Qui a dit : « Dans tout théoricien de la démonstration il y a un bureaucrate qui sommeille. »?

Ainsi, λxx représente bien la fonction identique : $\lambda xx(u) = u$. L'équation (4) peut être utilisée comme algorithme de calcul, et le théorème de Church-Rosser assure que ce type de calcul est confluent, i.e. ne donne pas de résultats contradictoires.

Cela étant, la proximité à la théorie des ensembles suggère de traduire le paradoxe de Russell

$$\{x; \neg x \in x\} \in \{x; \neg x \in x\} \leftrightarrow \neg(\{x; \neg x \in x\} \in \{x; \neg x \in x\}) \quad (5)$$

ce qui donne

$$\lambda x M(x(x))(\lambda x M(x(x))) = M(\lambda x M(x(x))(\lambda x M(x(x)))) \quad (6)$$

soit une équation $a = M(a)$. A propos qu'est-ce donc que ce terme M ? C'est un terme arbitraire qu'on a pris pour traduire la négation « $\neg$ ». Dans les deux cas, nous avons un point fixe, ici pour la négation, là pour un M arbitraire ; or si la logique classique n'admet pas de point fixe pour la négation, l'algorithmique s'accommode aisément des points fixes, tout au plus le calcul de $M(a)$ divergera... De fait le lambda-calcul est non seulement valide du point de vue algorithmique, mais c'est aussi une algorithmique universelle : cela résulte de la possibilité de programmer par point fixe que nous venons d'établir, et aussi de la possibilité de représenter les entiers par des termes, typiquement représenter n par $\lambda y \lambda x y(\dots(y(x)\dots))$, i.e. $\lambda y \lambda x x$, $\lambda y \lambda x y(x)$, $\lambda y \lambda x y(y(x))$, ... On s'amusera à calculer $n(m)$ ⁴⁰ pour se convaincre de la puissance calculatoire diabolique du formalisme.

De même que Russell avait proposé en 1905 sa théorie des types pour remédier aux contradictions ensemblistes, on a proposé le typage pour remédier à la non-terminaison de l'algorithme de calcul : c'est tout simple, on fabrique des types à partir de types de base non spécifiés $\alpha, \beta, \gamma, \dots$ au moyen de la flèche $\sigma \Rightarrow \tau$ qui est le type des fonctions de σ dans τ . Dans le calcul typé, tous les termes, à commencer par les variables, reçoivent un type bien précis. L'application $t(u)$ n'est possible que quand t a reçu un type $\sigma \Rightarrow \tau$ et u le type σ , auquel cas $t(u)$ est de type τ , l'abstraction $\lambda x t$ reçoit le type $\sigma \Rightarrow \tau$, où σ est le type de x et τ celui de t ; de plus l'équation (4) est compatible avec le typage. Les entiers n peuvent recevoir tous les types $(\sigma \Rightarrow \sigma) \Rightarrow (\sigma \Rightarrow \sigma)$, que nous noterons $\iota(\sigma)$. Dans le calcul typé les calculs ne divergent pas, ce qui a comme corollaire immédiat (encore la diagonalisation !) qu'il ne s'agit pas d'une algorithmique universelle. Il est d'ailleurs facile de trouver une fonction non représentable : si la fonction m^n est bien typable (m de type $\iota(\sigma)$, n de type $\iota(\sigma \Rightarrow \sigma)$), la fonction n^n ne l'est pas.

Le système $\mathbb{F}$ est aussi un lambda-calcul typé convergent et qui ne peut donc pas être une algorithmique universelle. Cela dit il s'agit d'une restriction sans portée pratique, vu que tout algorithme dont on peut établir la terminaison en

40. réponse : m^n

arithmétique du second ordre (en pratique beaucoup plus que les mathématiques courantes) y est représentable : les seules fonctions non représentables sont celles qui sont construites... à cet effet par diagonalisation. Le système $\mathbb{F}$ admet une quantification sur les types $\Lambda\alpha\sigma$, où la variable α est muette. On peut maintenant abstraire un terme t de type σ par rapport à α , de façon à obtenir $\Lambda\alpha t$ de type $\Lambda\alpha\sigma$ ou appliquer un terme t de type $\Lambda\alpha\sigma$ à un type τ , de façon à obtenir $t\{\tau\}$ de type $\sigma[\tau/\alpha]$; l'équation

$$\Lambda\alpha t\{\tau\} = t[\tau/\alpha] \quad (7)$$

démontre strictement (4). Le Λ permet le polymorphisme, i.e. la possibilité de jouer sur les types : ainsi on pourra attribuer aux entiers le type $\Lambda\alpha i(\alpha)$, et $\Lambda\alpha n\{\alpha \Rightarrow \alpha\}(n\{\alpha\})$ permet de typer n^n et en pratique tous les algorithmes fonctionnels qui convergent.

Une expérience étrange est d'essayer de confondre la flèche $\Rightarrow$ et l'implication logique, le Λ et un quantificateur universel. Un type devient une formule, et un objet de ce type devient... une notation fonctionnelle pour une démonstration de ce type, vu comme formule. Ainsi le terme λxx de type $\sigma \Rightarrow \sigma$ (fonction identique) peut être vu comme une démonstration de la tautologie $\sigma \Rightarrow \sigma$. C'est l'*isomorphisme de Curry-Howard* entre lambda-calcul typé et déduction naturelle que nous venons de présenter sommairement. Comme il s'agit d'un véritable isomorphisme (avec transfert de structures), on peut faire la complète économie de la déduction naturelle, que nous verrons à nouveau passer à l'horizon du calcul des séquents, voir annexe B. Le fait que la déduction naturelle soit entre autres une version confluyente des séquents établit un très fort lien entre les annexes A et B, qui au fond parlent du même objet sous des formes apparemment très dissemblables. La tendance moderne est évidemment à ne voir là qu'une unique structure. Ainsi les *réseaux de démonstration* de la logique linéaire ne sont-ils plus que des graphes reliant des formules, soumis à des critères purement géométriques (e.g. absence de cycles d'une certaine forme) : il n'y a plus que l'objet, qui peut se voir comme une fonction, comme une démonstration, mais encore comme une interaction, une stratégie etc. Cette insistance récente sur un objet mathématique unique de la logique n'a été possible que par l'adéquation entre les contenus des annexes A et B, développées par des traditions distinctes, voire antagonistes : l'isomorphisme de Curry-Howard qui l'établit ne pouvait être un accident.

B Le calcul des séquents

Avant Gentzen, la logique s'écrit dans des formalismes « à la Hilbert », i.e. des axiomes et des règles. Les propriétés de ces systèmes sont désespérantes, témoins les deux seuls exercices possibles :

1. Démontrer l'équivalence de l'axiomatisation de M. Machin avec celle de M. Système.
2. Fabriquer un système avec un nombre minimal d'axiomes et de règles.

En particulier les systèmes à la Hilbert sont incapables d'expliquer ce qui différencie la logique (sans contenu spécifique) de systèmes formels adaptés à des situations particulières (e.g. la formalisation des algèbres de Lie).

En pratique, la règle essentielle est le *Modus Ponens* « De A et de $A \Rightarrow B$, déduire B », une règle aux très mauvaises propriétés : une démonstration de B se termine par un *Modus Ponens*, et la prémisse A n'a strictement rien à voir *a priori* avec B .

Techniquement parlant, Gentzen introduit des *séquents*, i.e. des expressions $\Gamma \vdash \Delta$ où $\Gamma (= A_1, \dots, A_n)$ et $\Delta (= B_1, \dots, B_m)$ sont des suites finies de formules. La signification intuitive de $\Gamma \vdash \Delta$ est que

$$A_1 \text{ et } \dots \text{ et } A_n \text{ impliquent } B_1 \text{ ou } \dots \text{ ou } B_m$$

Le calcul est divisé en trois groupes :

Groupe identité

Il contient l'axiome d'identité et la règle de coupure :

$$\frac{}{A \vdash A} \quad (id.) \qquad \frac{\Gamma \vdash \Delta, A \quad \Lambda, A \vdash \Pi}{\Gamma, \Lambda \vdash \Delta, \Pi} \quad (coup.)$$

Les deux règles mettent en relation des occurrences de la même formule, de part et d'autre du « portillon » $\vdash$. Ces règles sont *génériques*, puisque A est arbitraire ; elles expriment, de façons duales l'identité de A avec lui-même « A est A et réciproquement »⁴¹. La formulation est symétrique par rapport à la gauche et la droite. L'axiome d'identité n'est qu'une façon d'énoncer la tautologie $A \Rightarrow A$, tandis que la coupure n'est qu'une forme sophistiquée de *Modus Ponens* : témoin le cas particulier

$$\frac{\vdash A \quad A \vdash B}{\vdash B} \quad (coupure)$$

Le coup de génie de Gentzen c'est de dissocier deux utilisations du *Modus Ponens* : l'utilisation propre, qui correspond aux nécessités déductives est représentée par la coupure, et la gestion du symbole logique « $\Rightarrow$ » devient la règle logique gauche de l'implication.

Groupe structurel

Les règles structurelles permettent de manipuler de part et d'autre du symbole $\vdash$;

41. L'axiome dit qu'un A gauche est plus fort qu'un A droit, tandis que la règle dit qu'un A droit est plus fort qu'un A gauche.

ces manipulations sont génériques, i.e. ne supposent rien sur les formules ; elles ont aussi symétries par rapport à la gauche et la droite.

1. L'*échange* exprime la commutativité de la logique ;

$$\frac{\Gamma \vdash \Delta}{\sigma(\Gamma) \vdash \tau(\Delta)}$$

où σ, τ sont des permutations.

2. L'*affaiblissement* c'est le principe qu'on n'est pas tenu d'utiliser toutes ses hypothèses (principe de l'implication matérielle).

$$\frac{\Gamma \vdash \Delta}{\Gamma \vdash A, \Delta} \qquad \frac{\Gamma \vdash \Delta}{\Gamma, A \vdash \Delta}$$

3. La *contraction* énonce la possibilité de réutiliser une hypothèse (ce qui se démontre avec deux hypothèses A peut se démontrer avec une seule), c'est l'idempotence des opérations logiques, conjonction et disjonction.

$$\frac{\Gamma \vdash A, A, \Delta}{\Gamma \vdash A, \Delta} \qquad \frac{\Gamma, A, A \vdash \Delta}{\Gamma, A \vdash \Delta}$$

Ces règles réussissent à exprimer les propriétés de la conjonction et de la disjonction sans en introduire les symboles.

Groupe logique

Les règles *logiques* ne sont pas génériques, vu qu'elles varient suivant les formules : ici le symbole externe est essentiel. Ainsi la conjonction admet les règles :

$$\frac{\Gamma \vdash \Delta, A \quad \Gamma \vdash \Delta, B}{\Gamma \vdash \Delta, A \wedge B} \quad (\vdash \wedge) \qquad \frac{\Gamma, A \vdash \Delta}{\Gamma, A \wedge B \vdash \Delta} \quad (l\wedge \vdash)$$

$$\frac{\Gamma, B \vdash \Delta}{\Gamma, A \wedge B \vdash \Delta} \quad (r\wedge \vdash)$$

Une règle binaire introduit la conjonction à droite, deux règles unaires l'introduisent à gauche. C'est ici qu'est cachée la symétrie profonde du système : il y a orthogonalité (dans un sens que seule la logique linéaire a pu exprimer parfaitement) entre les règles gauches et droites : c'est le sens profond de l'élimination des coupures (voir plus bas). Il y a aussi une symétrie de nature plus globale : les règles de la négation

$$\frac{\Gamma, A \vdash \Delta}{\Gamma \vdash \Delta, \neg A} \quad (\vdash \neg) \qquad \frac{\Gamma \vdash \Delta, A}{\Gamma, \neg A \vdash \Delta} \quad (\neg \vdash)$$

permettent à une formule de franchir le « portillon » $\vdash$, autrement dit la négation exprime l'échange gauche/droite. Et donc, si on applique informellement le principe de De Morgan qui nous dit comment nier une conjonction, $\neg(A \wedge B) = \neg A \vee \neg B$, on peut espérer que les règles de la disjonction sont l'image miroir de celle de la conjonction... et c'est bien le cas :

$$\begin{array}{c} \frac{\Gamma \vdash \Delta, A}{\Gamma \vdash \Delta, A \vee B} \quad (\vdash l\vee) \\ \frac{\Gamma \vdash \Delta, B}{\Gamma \vdash \Delta, A \vee B} \quad (\vdash r\vee) \end{array} \qquad \frac{\Gamma, A \vdash \Delta \quad \Gamma, B \vdash \Delta}{\Gamma, A \vee B \vdash \Delta} \quad (\vee \vdash)$$

La même symétrie échange les quantificateurs $\forall$ et $\exists$; seul connecteur autodual, $\neg$ a des règles symétriques ; seule l'implication fait cavalier seul, car son dual $\neg A \wedge B$ est dépourvu d'intérêt. Il est à noter que le choix des règles logiques est unique à quelques variantes près (équivalentes modulo règles structurelles).

Il est facile de voir que le calcul des séquents est une formulation parmi d'autres de la logique, d'ailleurs assez lourde. La seule excuse pour son introduction est la fameux *Hauptsatz*, i.e. l'élimination des coupures.

B.1 L'élimination des coupures

Un peu de terminologie :

- Dans toutes les règles structurelles ⁴² et logiques, il y a une conclusion « distinguée », e.g. la formule $A \wedge B$ dans la conclusion $\Gamma \vdash \Delta, A \wedge B$ de la règle $(\vdash \wedge)$; on l'appelle la *formule principale* de la règle ; les autres formules de la conclusion forment le *contexte*.
- La A qui apparaît deux fois dans la règle de coupure

$$\frac{\Gamma \vdash \Delta, A \quad \Lambda, A \vdash \Pi}{\Gamma, \Lambda \vdash \Delta, \Pi} \quad (coup.)$$

s'appelle la *formule de coupure*.

L'élimination des coupures résulte de la combinaison de plusieurs cas :

- Une coupure de formule de coupure A et telle que les deux occurrences de A sont les formules principales de règles logiques peut facilement être remplacée par des coupures plus simples : c'est ce qu'on appelle les cas-clef ;
- La situation devient nettement plus embrouillée quand A est conséquence d'un affaiblissement ou d'une contraction (technique des coupures croisées) ;

⁴². Ça ne vaut pas pour l'échange, qui est une règle à part, mais dans la pratique on travaille modulo permutation, i.e. sans échange.

- Plus généralement, quand A n'est pas formule principale, des commutations de règles doivent être effectuées ;
- Les cas précédents éliminent une coupure donnée en faveur de (beaucoup de) coupures sur des formules plus simples, et la taille de la démonstration s'accroît en général ; une itération intelligente de ces cas précédents finit par détruire toutes les coupures.

Les cas-clef

Prenons une coupure de formule de coupure A , dont les deux occurrences sont des formules principales de règles logiques : typiquement (avec $A = B \wedge C$)

$$\frac{\frac{\Gamma \vdash \Delta, B \quad \Gamma \vdash \Delta, C}{\Gamma \vdash \Delta, B \wedge C} \quad (\vdash \wedge) \quad \frac{\Lambda, C \vdash \Pi}{\Lambda, B \wedge C \vdash \Pi} \quad (r\wedge \vdash)}{\Gamma, \Lambda \vdash \Delta, \Pi} \quad (coup.)$$

On peut le remplacer par

$$\frac{\Gamma \vdash \Delta, C \quad \Lambda, C \vdash \Pi}{\Gamma, \Lambda \vdash \Delta, \Pi} \quad (coup.)$$

où la formule de coupure est C , donc plus simple. La possibilité de simplifier les coupures-clef est la symétrie la plus profonde du système, puisqu'elle suppose un certain équilibre entre les règles gauches et droites d'un même connecteur.

Cas des règles structurelles

La situation devient plus complexe quand une des occurrences de A est formule principale d'un affaiblissement ou d'une contraction. La tentation est de simplifier une telle coupure par un effacement ou une duplication, ainsi

$$\frac{\Gamma \vdash \Delta, A \quad \frac{\Lambda \vdash \Pi}{\Lambda, A \vdash \Pi} \quad (aff. \vdash)}{\Gamma, \Lambda \vdash \Delta, \Pi} \quad (coup.)$$

peut être remplacée par des affaiblissements itérés de $\Lambda \vdash \Pi$ (et ainsi le morceau de démonstration se terminant avec $\Gamma \vdash \Delta, A$ est effacé) ; de même

$$\frac{\Gamma \vdash \Delta, A \quad \frac{\Lambda, A, A \vdash \Pi}{\Lambda, A \vdash \Pi} \quad (contr. \vdash)}{\Gamma, \Lambda \vdash \Delta, \Pi} \quad (coup.)$$

peut être remplacée par deux coupures :

$$\frac{\Gamma \vdash \Delta, A \quad \frac{\Gamma \vdash \Delta, A \quad \Lambda, A, A \vdash \Pi}{\Gamma, \Lambda, A \vdash \Delta, \Pi} \quad (coup.)}{\Gamma, \Gamma, \Lambda \vdash \Delta, \Delta, \Pi} \quad (coup.)$$

suivies de contractions. La contraction induit donc une duplication du morceau se terminant avec $\Gamma \vdash \Delta, A$. Ce qui fait que la procédure ne converge pas si les deux occurrences of A sont obtenues par contraction, chacune dupliquant l'autre. Il y a plusieurs moyens de s'en sortir, en général par la technique des coupures croisées (trop technique pour cette introduction), ou au moyen de restrictions qui empêchent ce cas de se produire : par exemple en logique intuitionniste où il n'y a pas de règles structurelles droites.

Les commutations de règles

En général il n'y pas de raison que les occurrences de A dans une coupure soient des formules principales. Mais dans ce cas on peut toujours faire « remonter » la coupure ; ainsi une coupure

$$\frac{\Gamma \vdash \Delta, A \quad \frac{\Lambda, A, B \vdash \Pi}{\Lambda, A, B \wedge C \vdash \Pi} (\wedge \vdash)}{\Gamma, \Lambda, B \wedge C \vdash \Delta, \Pi} (coup.)$$

peut être remplacée par

$$\frac{\frac{\Gamma \vdash \Delta, A \quad \Lambda, A, B \vdash \Pi}{\Lambda, B \vdash \Pi} (coup.)}{\Gamma, \Lambda, B \wedge C \vdash \Delta, \Pi} (\wedge \vdash)$$

La commutation modifie le contexte d'application des règles : c'est l'origine des contextes dans la formulation des règles.

L'algorithme

En itérant de la bonne façon, les coupures finissent par disparaître : les cas de disparition ne sont que deux

- Dans le cas d'un affaiblissement, comme nous l'avons déjà vu.
- Quand une des deux prémisses est un axiome d'identité.

La taille des démonstrations a tendance à augmenter (reponsable : la contraction), par contre les formules de coupure se simplifient progressivement, puisqu'une coupure sur A est remplacée par des coupures sur des sous-formules de A .

B.2 La propriété de la sous-formule

Etant donnée une formule A , dont nous soupçonnons qu'elle est prouvable, dans quelle mesure peut-on synthétiser une éventuelle démonstration ? Formulons la question dans le cadre du calcul des séquents, en cherchant à démontrer $\vdash A$. En présence de la coupure, la question est sans espoir : dans une « vraie » démonstration, le théorème A est forcément construit comme un enchaînement de lemmes,

assemblés par le *Modus Ponens*, i.e. la coupure, et ces lemmes (qui contiennent les idées) sont imprévisibles. En termes techniques, à partir de la conclusion d'une coupure, il n'y a aucun moyen de prédire la formule de coupure. La question devient beaucoup plus raisonnable si on se restreint à des démonstrations sans coupures (ce qui a un sens du fait du *Hauptsatz*) : en effet, on vérifie facilement qu'une pré-misse $\Gamma' \vdash \Delta'$ de n'importe quelle règle autre que la coupure, et de conclusion $\Gamma \vdash \Delta$, est dans une certaine mesure « plus simple » que $\Gamma \vdash \Delta$, e.g. :

- Dans le cas d'une règle structurelle, $\Gamma' \vdash \Delta'$ est obtenu à partir de $\Gamma \vdash \Delta$ au moyen d'une permutation, d'un effacement ou d'une duplication ;
- Dans le cas d'une règle logique propositionnelle, $\Gamma' \vdash \Delta'$ est fait de formules déjà présentes dans $\Gamma \vdash \Delta$, sauf pour certaines formules qui sont des sous-formules d'une formule de la conclusion (e.g. $A \wedge B$ est dans $\Gamma \vdash \Delta$, et B est dans $\Gamma' \vdash \Delta'$) ;
- Le cas de la quantification est similaire au cas propositionnel, mais on doit tenir compte de la possibilité de substitution (e.g. $\forall x A$ est dans $\Gamma \vdash \Delta$, et $A[t/x]$ est dans $\Gamma' \vdash \Delta'$).

La notion de sous-formule « à la Gentzen » combine la notion de sous-formule et celle de substitution. Elle permet d'énoncer la propriété : pour démontrer $\vdash A$, on peut se restreindre à des séquents formés de sous-formules de $\vdash A$.

La propriété de la sous-formule restreint tellement l'espace de recherche, qu'il s'en faut d'un cheveu que la recherche de démonstration ne soit décidable. En tout cas elle induit des algorithmes de démonstration automatique non désespérés, et en particulier (la partie logique de) PROLOG.

B.3 Le cas intuitionniste

Le calcul intuitionniste est obtenu en se restreignant à des séquents $\Gamma \vdash A$. La présence d'une seule formule à droite rend caduques les règles structurelles droites, et simplifie radicalement le problème de l'élimination des coupures (plus de conflit contraction/contraction). Mais la première conséquence spectaculaire est la *propriété de la disjonction* ⁴³ : si $\vdash A \vee B$ est démontrable, alors $\vdash A$ est démontrable ou $\vdash B$ est démontrable. En effet, une fois les coupures éliminées, il ne reste que les deux règles logiques droites pour la disjonction ⁴⁴. C'est la base technique de l'explicitation et de l'importance algorithmique du théorème de Gentzen.

Il se pose encore une question : je me donne une démonstration Π avec coupures de $\vdash A \vee B$, et le *Hauptsatz* induit un choix gauche/droite entre $\vdash A$ et $\vdash B$; ce choix est-il prédéterminé par Π , où peut-il être influencé par les particularités de mon algorithme d'élimination ? La réponse est surprenante et sans recours : le

43. Il y a une propriété similaire pour le quantificateur existentiel.

44. En logique classique, la dernière règle est pratiquement toujours une contraction.

Hauptsatz est profondément déterministe, en particulier Π contient implicitement soit une démonstration de $\vdash A$, soit une démonstration de $\vdash B$, mais pas les deux. Ce résultat fondamental est obtenu en remplaçant le calcul des séquents par la *déduction naturelle*, et en prouvant alors un résultat de *confluence* pour cette variante du calcul des séquents, i.e. que l'algorithme d'élimination est déterministe. Puisque la déduction naturelle est isomorphe au lambda-calcul typé, expliquons l'idée dans ce cadre : une démonstration d'un séquent, disons $A, B \vdash C$ est interprétée par un terme t de type C dépendant de deux variables x, y de types A, B (informellement, par une fonction de deux arguments envoyant $A \times B$ dans C). La confluence n'est rien d'autre que la propriété de Church-Rosser du lambda-calcul. C'est ici que les deux sentiers se rejoignent.

B.4 La linéarité

La logique intuitionniste n'a pas de négation involutive, i.e. le raisonnement par l'absurde n'est pas valide dans ce cas. En fait c'est parce que le seul sens possible de la négation est l'échange gauche/droite dans un séquent et que la gauche et la droite ne sont pas équivalents du point de vue intuitionniste. Il est cependant possible de concilier la symétrie gauche/droite (i.e. la négation involutive) et la propriété de la disjonction en interdisant les règles d'affaiblissement et de contraction (seule leur absence à droite a été utilisée). L'interprétation fonctionnelle du calcul des séquents fait alors apparaître des fonctions linéaires, d'où le nom de *logique linéaire*. Mais il s'agit avant tout d'une symétrisation de la logique intuitionniste : l'affaiblissement et la contraction peuvent se voir comme la cointé et la comultiplication d'un comonoïde commutatif $!A$ (genre « algèbre tensorielle symétrique »), et donc il n'y a pas de perte d'expressivité.

En logique linéaire, on perd la directionalité, i.e. $A, B \vdash C$ peut aussi se voir comme une fonction bilinéaire de $C^\perp \otimes B$ dans $A^\perp$. Autant dire qu'on a dépassé le cadre fonctionnel. Quel est donc ce nouveau cadre ? Sans rentrer dans les détails, e.g. la géométrie de l'interaction, disons que les opérations logiques deviennent des opérations géométriques. Par exemple la négation est vraiment l'échange gauche/droite, l'affaiblissement la destruction physique, la contraction la duplication physique, la conjonction multiplicative $\otimes$ le multiplexage de canaux... quant à l'axiome d'identité il devient une rallonge entre deux prises complémentaires tandis que la coupure est le branchement de deux telles prises.

Bibliographie sélective

Le grand absent de cette bibliographie, c'est Kreisel : on ne peut hélas guère recommander des articles écrits dans un style King's College saturé d'allusions déjà peu compréhensibles à l'époque. On a quand même inclus l'article [27], un peu plus lisible que la moyenne.

J'ai indiqué un maximum de livres, quand c'était possible. En premier lieu, le livre de Jean van Heijenoort [39], qui traduit les textes fondateurs de la logique, en particulier le texte de Hilbert sur l'infini (1925) qui contient en outre une prétendue démonstration de l'hypothèse du continu, et le texte original de Gödel de 1931. On retrouve ce texte en traduction française dans le petit livre du Seuil [20], suivi d'un commentaire de l'auteur de cet article, avec entre les deux un texte de vulgarisation des philosophes Nagel et Newmann qui fait fort dans le réductionnisme Jivaro. Les articles de Gentzen ont été publiés en volume (non-disponible en ce moment) [14]. On trouve les œuvres logiques complètes de Herbrand en traduction anglaise [22]. La *Beweistheorie* de Schütte (1960) est disponible en traduction anglaise [33]. Si on s'intéresse à cette direction, il est possible de consulter [3]. La théorie de la démonstration classique est aussi présentée par Takeuti [36] et dans mon livre [17], dont la seconde partie, qui présente les dilatateurs sera peut-être publiée un jour. Pour l'intuitionnisme, on pourra consulter le pavé de Troelstra et van Dalen [37]. Malheureusement la thèse de Prawitz [32] n'est guère disponible, et une réédition serait la bienvenue. Le livre de Barendregt [1] est la référence sur le lambda-calcul. Le petit livre [18] est particulièrement léger et couvre de nombreux aspects de théorie de la démonstration sans se noyer dans les pinaillages. Le livre de Krivine [29], qui couvre des sujets voisins, a de plus le bon goût d'être en Français. Pour la logique linéaire, on pourra consulter le livre [19], qui contient des articles de recherche tout comme des exposés introductifs.

Références

- [1] H. Barendregt. **The lambda-calculus, its syntax and semantics**. *North Holland*, Amsterdam, 1984.
- [2] J. Bridge. **A simplification of the Bachmann method for generating large countable ordinals**. *Journal of Symbolic Logic*, 40:171–185, 1975.
- [3] W. Buchholz, S. Feferman, W. Pohlers, and W. Sieg. **Iterated Inductive Definitions and Subsystems of Analysis: Recent Proof-theoretical studies**. 897. *Springer Verlag*, Berlin, 1981. ISBN 3-540-11170-0.
- [4] C. Burali-Forti. **Sulle classi ben ordinate**. *Rendiconti del Circolo matematico di Palermo*, 11:260, 1897.
- [5] C. Burali-Forti. **Una questione sui numeri transfiniti**. *Rendiconti del Circolo matematico di Palermo*, 11:154–164, 1897.

- [6] S. R. Buss. **Bounded arithmetic**. *Bibliopolis*, Napoli, 1986. ISBN 88-7088-150-4.
- [7] T. Coquand and G. Huet. **Constructions : a higher order proof system for mechanizing mathematics**. In *EUROCAL'85*, Berlin, 1985. *Springer Verlag*. LNCS 203.
- [8] S. Feferman. **A language and axioms for explicit mathematics**. In *Algebra and Logic*, pages 87–139, Berlin, 1975. *Springer Verlag*. LNM 450.
- [9] J. H. Gallier. **What is so special about Kruskal's theorem and the ordinal Γ_0 ?** *Annals of Pure and applied logic*, 53:199–260, 1991.
- [10] G. Gentzen. **Untersuchungen über das logische Schliessen**. *Mathematische Zeitschrift*, 39:176–210, 405–431, 1935.
- [11] G. Gentzen. **Die Widerspruchsfreiheit der reinen Zahlentheorie**. *Mathematische Annalen*, 112:493–565, 1936.
- [12] G. Gentzen. **Die Widerspruchsfreiheit der Stufenlogik**. *Mathematische Zeitschrift*, 41.3:357–366, 1936.
- [13] G. Gentzen. **Neue Fassung des Widerspruchsfreiheitsbeweises für die reine Zahlentheorie**. *Forschungen zur Logik und zur Grundlegung der exakten Wissenschaften, nouvelle série*, 4:19–44, 1938. Leipzig (Hirzel).
- [14] G. Gentzen. **The collected works of Gerhard Gentzen**, ed. Szabo. *North Holland*, Amsterdam, 1969. SBN 7204-2254-X.
- [15] J.-Y. Girard. **Π_2^1 -logic, part I: dilators**. *Annals of Mathematical Logic*, 21:75–219, 1981.
- [16] J.-Y. Girard. **The Ω -rule**. In *Proceedings of the International Congress of Mathematicians*, pages 307–321, Warszawa, 1984. PWN.
- [17] J.-Y. Girard. **Proof-theory and logical complexity I**. *Bibliopolis*, Napoli, 1987.
- [18] J.-Y. Girard, Y. Lafont, and P. Taylor. **Proofs and types**, volume 7 of *Cambridge tracts in theoretical computer science*. *Cambridge University Press*, Cambridge, 1990.
- [19] J.-Y. Girard, Y. Lafont, and L. Regnier, editors. **Advances in Linear Logic**, volume 222 of *London Mathematical Society Lecture Note Series*. *Cambridge University Press*, Cambridge, 1995. ISBN 0-521-55961-8.
- [20] K. Gödel, E. Nagel, J. R. Newmann, and J.-Y. Girard. **Le théorème de Gödel**. Sources du Savoir. *Le Seuil*, Paris, 1989. ISBN 2-02-010652-3.

- [21] J. Herbrand. **Sur le problème fondamental de la logique mathématique.** *Sprawozdania z posiedzeń Towarzystwa Naukowego Warszawskiego, Wydział III*, 24:12–56, 1931.
- [22] J. Herbrand. **Collected Works**, ed. Goldfarb. *Harvard University Press*, Cambridge, Massachusetts, 1971. SBN 674-80206-3.
- [23] D. Hilbert. **Über die Grundlagen der Logik und die Arithmetik.** *Verhandlungen des Dritten Internationalen Mathematiker-Kongresses in Heidelberg*, 1905.
- [24] D. Hilbert. **Über das Unendliche.** *Mathematische Annalen*, 95:161–190, 1926.
- [25] G. Kreisel. **On the interpretation of non-finitistic proofs I.** *Journal for Symbolic Logic*, 16:241–267, 1951.
- [26] G. Kreisel. **On the interpretation of non-finitistic proofs II.** *Journal for Symbolic Logic*, 17:43–58, 1952.
- [27] G. Kreisel. **A survey of proof-theory.** *Journal of Symbolic Logic*, 33:321–388, 1968.
- [28] G. Kreisel and A. Levy. **Reflection principles and their use for establishing the complexity of logical systems.** *Zeitschrift für Mathematische Logik*, 14:97–142, 1968.
- [29] J.-L. Krivine. **Lambda-calcul, types et modèles.** *Masson*, Paris, 1990. ISBN 2-225-82091-0.
- [30] J.-L. Krivine. **Classical logic, storage operators and second order lambda-calculs.** *Annals of Pure and Applied Logic*, 68:53–78, 1994.
- [31] P. Martin-Löf. **Intuitionistic Type Theory.** *Bibliopolis*, Napoli, 1984. ISBN 88-7088-105-9.
- [32] D. Prawitz. **Natural Deduction.** *Almqvist & Wiksell*, Stockholm, 1965.
- [33] K. Schütte. **Proof-theory.** *Springer Verlag*, Berlin, 1977.
- [34] D. Scott. **Data types as lattices.** *SIAM Journal of Computing*, 5:522–587, 1976.
- [35] G. Takeuti. **Consistency proofs for subsystems of classical analysis.** *Annals of Mathematics*, 86:299–348, 1967.
- [36] G. Takeuti. **Proof-theory.** *North Holland*, Amsterdam, 1975. ISBN 0-444-10492-5.

- [37] A. S. Troelstra and D. van Dalen. **Constructivism in Mathematics, vols 1 & 2.** *North Holland*, Amsterdam, 1988. ISBN 0-444-70266-0, 0-444-70358-6.
- [38] D. van Dalen. **The War of the Frogs and the Mice, or the Crisis of the Mathematische Annalen.** *Mathematical Intelligencer*, 12:17–31, 1990.
- [39] J. van Heijenoort. **From Frege to Gödel.** *Harvard University Press*, Cambridge, Massachusetts, 1967. SBN 674-32450-1.