

Sur le théorème de complétude de la logique propositionnelle

1 Introduction

On se propose de produire une preuve indépendante du théorème de complétude de la logique propositionnelle, qui s'énonce essentiellement "Une formule propositionnelle vraie est prouvable". On prouvera au passage l'équivalence de cet énoncé avec un énoncé qui est une conséquence de l'axiome du choix mais pas des autres axiomes de ZF seuls.

2 La logique propositionnelle classique

2.1 Définition, contexte

Avant d'énoncer et de prouver le théorème en question, il faut voir de quoi il parle.

La logique propositionnelle classique est une logique assez "primitive", au sens où elle est peu expressive et ne permet pas de gérer des affirmations comme "Pour tout réel x , il existe un réel y tel que ...". Elle ne permet que des affirmations comme "Si (s'il pleut alors j'apporte mon parapluie) alors (si je n'apporte pas mon parapluie alors il ne pleut pas)". Elle est tout de même intéressante, que ce soit pour elle-même en tant que système logique, ou pour l'étude d'autres logiques plus expressives comme la logique du premier ordre classique. Elle est de plus malgré tout expressive, et plus "forte" que d'autres, comme la logique intuitionniste.

Pour définir les formules propositionnelles on considère:

- Un ensemble infini $\mathcal{V}$ de variables (dans le cas de base, il sera dénombrable, et on notera les variables $A_1, \dots, A_n, \dots$);
- Des connecteurs binaires $\implies, \vee, \wedge$;
- Un connecteur unaire $\neg$;
- Des constantes $\perp, \top$

On aura aussi besoin de parenthèses, pour tout simplifier $), ($, qui seront des symboles. On suppose que tous ces symboles sont deux à deux distincts.

On pourrait faire une liste plus courte, mais tout serait moins lisible, moins compréhensible et surtout moins intuitif: ici on peut justifier et donner un sens "naturel" à chacun des connecteurs. Les variables représentent des données qui peuvent varier, comme par exemple "il pleut" (l'interprétation de cette variable comme "vraie" ou "fausse" dépend de là où vous êtes, quand vous lisez cet article, et surtout de quelle définition vous donnez de "il pleut").

Les connecteurs représentent le moyen de construire des phrases plus complexes à partir de ces variables comme "il pleut et je prends mon parapluie", qu'on écrirait $A_1 \wedge A_2$ si A_1 représente "il pleut" et A_2 "je prends mon parapluie". $\wedge$ correspond à la conjonction, donc au "et" usuel, $\vee$ à la disjonction, donc au "ou" non exclusif (si je dis "il pleut ou je mange", il n'est pas exclu que je mange et qu'il pleuve), $\implies$ à l'implication: $P \implies Q$ est vraie si, si P est vraie alors Q aussi; $\neg$ correspond à la négation usuelle: la négation de "il pleut" est "non (il pleut)" ou plus naturellement "il ne pleut pas". $\perp$ et $\top$ sont des aides pour faciliter la construction de ce qui suit. $\perp$ est une constante qui est toujours fausse, c'est "le faux"; au contraire $\top$ est une constante toujours vraie, c'est "le vrai".

Etant donné un ensemble X on appelle "monoïde libre sur X " et on note X^* l'ensemble des

chaînes de caractères de X . Par exemple si $X = \{0, 1\}$, $011001 \in X^*$. Si $x \neq 0, 1$, alors $1x0 \notin X^*$. Plus formellement, $X^* = \bigcup_{n \in \mathbb{N}} X^n$. Noter qu'on autorise $n = 0$, sachant que $X^0 = \{\emptyset\}$, cela revient

à dire qu'on autorise la chaîne vide, qu'on note dans ce contexte ϵ . On a une opération naturelle de concaténation sur X^* . On identifie X et X^1 de la manière évidente, c'est-à-dire que l'élément $x \in X$ correspond à l'application $0 \mapsto x$.

L'ensemble des formules propositionnelles est un sous-ensemble de F^* pour un F approprié. Plus précisément, soit $F = \mathcal{V} \cup \{ \implies, \wedge, \vee \} \cup \{ \neg \} \cup \{ \top, \perp \} \cup \{ (,) \}$.

On dit qu'un sous-ensemble S de F^* a la propriété P si :

- $\mathcal{V} \subset S$
- Si $P, Q \in S$, alors $(P \wedge Q), (P \vee Q), (P \implies Q) \in S$
- Si $P \in S$ alors $\neg P \in S$
- $\perp, \top \in S$

On note ça $P(S)$. Il est clair que F^* a la propriété P . Soit alors $\mathcal{F} := \bigcap \{ S \subset F^* \mid P(S) \}$. Il est clair que $\mathcal{F}$ a la propriété P (si vous n'êtes pas convaincu, vérifiez point par point, vous verrez que ça n'a rien de sorcier). On appelle $\mathcal{F}$ l'ensemble des formules propositionnelles. Il est de plus clair que c'est le plus petit sous-ensemble de F^* à l'avoir: il est par définition inclus dans tous ceux qui l'ont. Ce fait nous permettra de mettre en place une technique de preuve par induction efficace. On la décrit de manière générale:

Théorème (Preuve par induction structurelle) Soit Q une propriété que les éléments de F^* peuvent avoir ou ne pas avoir. Supposons que Q soit vérifiée par les éléments de $\mathcal{V}$, par $\perp, \top$. Supposons de plus que si A, B vérifient Q alors $(A \implies B), (A \wedge B), (A \vee B)$ la vérifient aussi, et finalement supposons que si A vérifie Q , alors $\neg A$ aussi. Alors tous les éléments de $\mathcal{F}$ la vérifient aussi.

Preuve du théorème Considérons l'ensemble $T = \{x \in F^* \mid Q(x)\}$. Par les hypothèses, $P(T)$. D'après ce qui précède on en déduit $\mathcal{F} \subset T$, i.e. tous les éléments de $\mathcal{F}$ vérifient Q .

Cela permet de donner une autre caractérisation de $\mathcal{F}$. En effet, considérons $A_0 = \mathcal{V} \cup \{ \perp, \top \}$. On a envie de dire qu'il suffit d'ajouter les implications, conjonctions et disjonctions entre les éléments de A_0 , puis les négations de ces éléments. On obtient alors A_1 . Mais A_1 n'est lui-même pas stable par ces opérations, il faut donc répéter le processus et obtenir A_2 . On continue ainsi, en définissant $A_{n+1} := A_n \cup \{ \neg P \mid P \in A_n \} \cup \{ (P \implies Q), (P \vee Q), (P \wedge Q) \mid P, Q \in A_n \}$. On note alors $\mathcal{F}' := \bigcup_{n \in \mathbb{N}} A_n$. Soit la propriété $Q(x) \iff x \in \mathcal{F}'$. Il est clair que Q vérifie, par construction,

les hypothèses du théorème d'induction structurelle: donc $\mathcal{F} \subset \mathcal{F}'$. Pour obtenir la réciproque, on procède par récurrence sur $\mathbb{N}$, en montrant qu'à chaque étape, $A_n \subset \mathcal{F}$. Ainsi $\mathcal{F} = \mathcal{F}'$. On a ainsi une description "par le bas" de $\mathcal{F}$.

Nous avons pour le moment créé des formules sans leur donner de sens: nous avons fait attention à la syntaxe, sans pour autant faire attention à la sémantique. C'est un peu comme si en regardant la phrase "Le chat pleut hier" on ne remarquait pas qu'elle était bizarre parce qu'elle vérifie bien la syntaxe sujet-verbe-complément. Mais avant de passer à la sémantique il faut vérifier tous les détails et vérifier des petites choses techniques. En voici quelques unes:

Théorème de lecture unique: Tout élément de $\mathcal{F}$ s'écrit de manière unique $\neg P, (P \wedge Q), (P \vee Q), (P \implies Q), A_i, \perp$ ou $\top$ où $i \in \mathbb{N}$, $P, Q \in \mathcal{F}$, chacune des options s'excluant mutuellement.

Preuve: On veut voir que la propriété annoncée par ce théorème vérifie les conditions du théorème d'induction structurelle, et de manière à pouvoir l'appliquer pour conclure. La seule difficulté est de pouvoir distinguer $(P \wedge Q)$ de $(P \vee Q)$ et $(P \implies Q)$ par exemple. Pour ça il s'agit de savoir récupérer P à partir de $(P \wedge Q)$ (par exemple). On laisse en exercice aux lectrices la vérification que l'algorithme suivant permet de retrouver P à partir de $(P \alpha Q)$, quel que soit α :

- Enlever la première parenthèse;
- S'il y a un $\neg$, tous les lire puis passer à la prochaine étape, sinon, passer à la prochaine étape;
- Si on lit une variable ou $\perp$ ou $\top$, on s'arrête, le compteur reste à 0. Sinon c'est qu'on lit une parenthèse (. Ajouter 1 au compteur;
- Tant que le compteur est > 0 , lire tous les caractères, en ajoutant 1 au compteur si on lit un

parenthèse (, enlevant 1 au compteur si on lit une parenthèse);

-Si le compteur est à 0, la liste des caractères lus est P .

Ainsi on peut retrouver P et donc α et donc Q , ce qui assure l'unicité à nouveau.

Ce théorème technique permet de définir les valuations, qui sont le coeur de la sémantique propositionnelle.

Définition (valuation): Une valuation δ est une application $\delta : \mathcal{V} \rightarrow \{0, 1\}$.

Le théorème de lecture unique permet d'assurer que la définition suivante est justifiée:

Reprenant les notations d'avant, si δ est une valuation, on définit son extension, ici notée $\bar{\delta}$ mais souvent notée δ par abus de notation, à l'ensemble des formules propositionnelles: $\bar{\delta}(A_i) := \delta(A_i)$ pour tout i , $\bar{\delta}(\perp) := 0$, $\bar{\delta}(\top) := 1$, $\bar{\delta}((P \wedge Q)) := \min(\bar{\delta}(P), \bar{\delta}(Q))$, $\bar{\delta}((P \vee Q)) := \max(\bar{\delta}(P), \bar{\delta}(Q))$, $\bar{\delta}(\neg P) := 1 - \bar{\delta}(P)$, et finalement $\bar{\delta}((P \implies Q)) := \max(1 - \bar{\delta}(P), \bar{\delta}(Q))$

Le théorème de lecture unique permet d'affirmer que $\bar{\delta}$ est définie de manière unique, et la description par le bas de $\mathcal{F}$ montre que $\bar{\delta}$ est bien définie (on peut choisir de le faire par récurrence par exemple) et qu'elle est définie sur tout $\mathcal{F}$.

Toutes ces définitions sont claires, sauf peut-être celle liée à $\implies$: il faut voir que l'extension de δ donne la même valeur à $(P \implies Q)$ qu'à $(\neg P \vee Q)$. C'est parce que l'interprétation mathématique de $P \implies Q$ fait que ces deux formules sont "équivalentes". Le raisonnement qui suit est intuitif et est là pour justifier cette identification: supposons $\neg P \vee Q$, on veut montrer $P \implies Q$. Il faut donc montrer que "si P , alors Q ". Mais si P est vraie, alors $\neg P$ ne l'est pas. Mais on sait que $\neg P$ est vraie, ou Q l'est, car $\neg P \vee Q$ est vraie. Mais puisque $\neg P$ ne l'est pas, c'est que Q doit l'être. On a donc montré que si P est vraie, Q l'est. Réciproquement, supposons $P \implies Q$ et montrons $\neg P \vee Q$. Si P n'est pas vraie, c'est que $\neg P$ l'est, donc $\neg P$ ou Q est vraie (puisque $\neg P$ l'est !). Si au contraire, P est vraie, alors puisque $P \implies Q$ l'est, c'est que Q l'est, et donc $\neg P \vee Q$ aussi (puisque Q est vraie !). On a donc montré intuitivement l'équivalence, ce qui justifie la définition ci-dessus. Remarquez que cette "preuve" (qu'on saura formaliser plus tard) n'est valable que parce qu'on a supposé que P est vraie, ou $\neg P$ l'est, ce qui est du ressort de la logique classique et pas d'autres logiques plus exotiques comme la logique intuitionniste par exemple.

La définition de valuation et d'extension de valuation (qu'on appelle aussi valuation, par abus) permet de donner une interprétation aux formules. Admettons qu'on a la formule "Si il pleut, je prends mon parapluie". On note P "il pleut" et Q "je prends mon parapluie". On veut savoir si en ce moment, $P \implies Q$ est vraie. Pour cela je regarde si il pleut (je donne une évaluation à P), et si je prends mon parapluie (j'en donne une à Q). J'associe les réponses à 0, 1 selon la réalité, et j'obtiens ainsi $\delta : \{P, Q\} \rightarrow \{0, 1\}$ (on se fiche ici des autres variables), et je peux alors calculer $\bar{\delta}((P \implies Q))$, qui me dit si je respecte ma règle ou non. Si la valuation étendue vaut ici 1 pour une certaine valuation δ , cela ne veut pas pour autant dire que $P \implies Q$ est vraie, simplement qu'en cette instance, elle l'est.

Définition: Soit $P \in \mathcal{F}$ et δ une valuation. On note $\delta \models P$ si et seulement si $\bar{\delta}(P) = 1$. Il faut lire quelque chose comme " δ satisfait P " ou "sous l'interprétation δ , P est vraie". Si $\Gamma \subset \mathcal{F}$, on note aussi $\delta \models \Gamma$ si et seulement si $\delta \models Q$ pour tout $Q \in \Gamma$.

On note finalement $\Gamma \models P$ si pour toute valuation δ telle que $\delta \models \Gamma$, on a aussi $\delta \models P$. On dit alors quelque chose comme " Γ entraîne P ", ou " P est conséquence sémantique de Γ ". En particulier, si aucune valuation δ ne vérifie $\delta \models \Gamma$, alors pour toute formule P , $\Gamma \models P$. On dit alors que Γ est inconsistante. Si au contraire il existe une telle valuation δ , on dit que Γ est consistante.

L'exercice suivant est laissé aux lectrices : Γ est inconsistante si et seulement si toute formule en est conséquence sémantique, si et seulement si $\perp$ en est conséquence sémantique, si et seulement si il existe une formule P telle que $\Gamma \models (P \wedge \neg P)$.

On laisse aussi en exercice les propriétés suivantes de base: pour une valuation δ , $\delta \models (P \wedge Q)$ si et seulement si $\delta \models P$ et $\delta \models Q$, $\delta \models (P \vee Q)$ si et seulement si $\delta \models P$ ou $\delta \models Q$; si $\delta \models P$ et $\delta \models (P \implies Q)$ alors $\delta \models Q$.

Définition : Une formule P est dite valide, ou universellement valide si pour toute valuation δ , $\delta \models P$.

le. a lectrice verra rapidement que P est valide si et seulement si $\emptyset \models P$, ce que l'on note aussi $\models P$.

le.a lectureurice peut déjà nous voir se rapprocher de l'énoncé du théorème de complétude, qu'on peut déjà entrevoir : "Une formule propositionnelle est valide si et seulement si elle est prouvable". Reste à voir ce que "prouvable" veut dire...

La relation $\models \subset \mathcal{P}(\mathcal{F}) \times \mathcal{F}$ est une relation dite de conséquence. En fait si on connaissait le théorème, on saurait que $\Gamma \models P$ revient à dire qu'on peut prouver P à partir des hypothèses Γ , ce qui explique le nom conséquence. Cette relation donne une opération $C : \mathcal{P}(\mathcal{F}) \rightarrow \mathcal{P}(\mathcal{F})$ avec $C(\Gamma) = \{P \mid \Gamma \models P\}$. Il est évident au vu des définitions que $\Gamma \subset C(\Gamma)$, que si $\Gamma \subset \Omega$, $C(\Gamma) \subset C(\Omega)$ et on vérifie (ce n'est pas compliqué mais ça l'est plus que les deux autres) que $C(C(\Gamma)) = C(\Gamma)$. En somme, c'est ce que l'on attend d'un "opérateur de conséquence".

En effet supposons que l'on ait $C : \mathcal{P}(\mathcal{F}) \rightarrow \mathcal{P}(\mathcal{F})$ qui représente "À Γ j'associe l'ensemble des conséquences de Γ ". On voudrait alors bien que $\Gamma \subset C(\Gamma)$: une formule est conséquence d'elle-même; que si $\Gamma \subset \Omega$, alors $C(\Gamma) \subset C(\Omega)$: en effet on ne supprime pas des conséquences en ajoutant des hypothèses; et finalement on voudrait que $C(C(\Gamma)) = C(\Gamma)$: une conséquence des conséquences est elle-même une conséquence. C'est en partant sur de telles considérations qu'on peut généraliser notre étude à des logiques bien plus générales et abstraites que la logique propositionnelle classique. Certains rajoutent aussi la condition qu'une conséquence de Γ soit en fait conséquence d'un sous-ensemble fini de Γ , ce à quoi on pourrait s'attendre. On verra grâce au théorème de complétude qu'ici c'est le cas; et on le redémontrera alternativement dans l'appendice. Mais on s'écarte du sujet, revenons-en au théorème de complétude.

Il nous faut donc définir ce qu'est la prouvabilité en logique propositionnelle. Pour prouver une formule, on peut partir d'hypothèses, pour arriver à une conclusion. Mais les hypothèses peuvent changer au cours de la preuve: imaginons que je veux prouver $P \implies Q$ à partir de Γ . Je veux à la fin obtenir que $P \implies Q$ découle de Γ . Mais durant la preuve, je peux être amené à dire "supposons P ", en déduire Q , et déduire de ça $P \implies Q$, en perdant l'hypothèse P . En plus de ça, il nous faut des règles d'inférences gouvernant tout cela. Voilà ce qu'on tente de formaliser dans la suite.

2.2 Un système de démonstration: les séquents

Un séquent est un couple (Γ, P) où $\Gamma \cup \{P\} \subset \mathcal{F}$. En général, on le note $\Gamma \vdash P$. Il faut le comprendre comme "De Γ , on peut inférer P ". En général dans le cours d'une preuve, on notera souvent $\Gamma, Q \vdash P$ pour le séquent $(\Gamma \cup \{Q\}, P)$. On définit alors la notion de démonstration.

Une démonstration est une suite finie D de séquents $\Gamma \vdash P$ vérifiant les conditions suivantes:

Si D_i est un séquent $\Gamma \vdash P$ alors l'une des possibilité suivantes doit être réalisée:

- $P \in \Gamma$ [les hypothèses sont prouvables]
- Il existe $j < i$, et une formule Q tels que D_j soit un séquent de la forme $\Gamma \vdash (P \wedge Q)$ ou $\Gamma \vdash (Q \wedge P)$ [si je sais prouver $P \wedge Q$ ou $Q \wedge P$, je sais prouver P]
- Il existe $j < i$, et des formules Q, R telles que $P = (Q \vee R)$, et tels que D_j soit un séquent de la forme $\Gamma \vdash Q$ ou $\Gamma \vdash R$ [si je sais prouver Q ou si je sais prouver R , je sais prouver $Q \vee R$]
- Il existe trois entiers $j, k, l < i$ et deux formules Q, R tels que D_j soit le séquent $\Gamma \vdash (Q \vee R)$, D_k le séquent $\Gamma, Q \vdash P$, et D_l le séquent $\Gamma, R \vdash P$ [Si avec Q je sais prouver P , avec R je sais prouver P , et si je sais prouver $Q \vee R$, alors je sais prouver P]
- $P = (Q \implies R)$ et il existe un entier $j < i$ tel que D_j soit le séquent $\Gamma, Q \vdash R$ [si à l'aide de Q je peux prouver R , je peux prouver $Q \implies R$]
- $P = (Q \wedge R)$ et il existe $j, k < i$ tels que D_j soit $\Gamma \vdash Q$, et D_k soit $\Gamma \vdash R$ [si je sais prouver Q et R , je sais prouver $Q \wedge R$]
- $P = \perp$ et il existe $j < i$, Q tels que D_j soit $\Gamma \vdash (Q \wedge \neg Q)$ [si je peux prouver qu'une certaine formule est vraie et fausse, je peux prouver le faux]

- $P = \top$ [je sais toujours prouver le vrai]
- Il existe $j < i$ tel que D_j soit $\Gamma \vdash \perp$. [si je sais prouver le faux, je sais prouver n'importe quoi]
- Il existe $j < i$ tel que D_j soit $\Gamma \vdash \neg\neg P$. [si je sais prouver que P n'est pas fausse, alors je sais prouver qu'elle est vraie]
- $P = \neg Q$ et il existe $j < i$ tel que D_j soit $\Gamma \vdash Q \implies \perp$ [si je sais prouver que Q implique le faux, je sais prouver que Q est fausse]
- Il existe $j, k < i$ et Q tels que D_j soit le séquent $\Gamma \vdash (Q \implies P)$ et D_k le séquent $\Gamma \vdash Q$ [Si je sais prouver que Q implique P et si je sais prouver Q , je sais prouver P]

Les commentaires écrits entre $\llbracket$ sont informels et sont là pour justifier intuitivement la règle d'inférence qu'on propose. On présentera une démonstration comme une suite de séquents numérotés, avec à chaque fois écrit "via telle règle, et les séquents tant et tant" à côté de sorte qu'un ordinateur pourrait vérifier algorithmiquement une démonstration. Les noms des règles d'inférences citées plus haut sont, dans l'ordre : Ax., $\wedge$ -elim, $\vee$ -intro, $\vee$ -elim, $\implies$ -intro, $\wedge$ -intro, $\perp$ -intro, $\top$ -intro, $\perp$ -elim, $\neg\neg$ -elim, $\neg$ -intro et MP. (modus ponens). Tout ceci est bien formel et abstrait, donnons donc des exemples.

Exemple basique: Prouvons que $\vdash P \implies P$ (abréviation de $\emptyset \vdash P \implies P$) est un séquent démontrable.

1. $P \vdash P$ (Ax.)
2. $\vdash P \implies P$ (de 1., $\implies$ -intro).

Autre exemple, plus profond que celui-là, car il donne un lien entre la double négation et le principe du tiers exclus:

1. $\neg(P \vee \neg P), \neg P \vdash \neg P$ (Ax.)
2. $\neg(P \vee \neg P), \neg P \vdash (P \vee \neg P)$ (de 1., $\vee$ -intro)
3. $\neg(P \vee \neg P), \neg P \vdash \neg(P \vee \neg P)$ (Ax.)
4. $\neg(P \vee \neg P), \neg P \vdash (P \vee \neg P) \wedge \neg(P \vee \neg P)$ (de 2. et 3., $\wedge$ -intro)
5. $\neg(P \vee \neg P), \neg P \vdash \perp$ (de 4., $\perp$ -intro)
6. $\neg(P \vee \neg P) \vdash (\neg P \implies \perp)$ (de 5., $\implies$ -intro)
7. $\neg(P \vee \neg P) \vdash \neg\neg P$ (de 6., $\neg$ -intro)
8. $\neg(P \vee \neg P) \vdash P$ (de 7., $\neg\neg$ -elim)

[Ici il est clair qu'on peut en quelques étapes refaire la même preuve et obtenir une démonstration $\neg(P \vee \neg P) \vdash \neg P$. Pour ne pas trop alourdir ce papier, je suppose qu'on l'a obtenu en le même nombre d'étapes]

16. $\neg(P \vee \neg P) \vdash \neg P$
17. $\neg(P \vee \neg P) \vdash (P \wedge \neg P)$ (de 8. et 16., $\wedge$ -intro)
18. $\neg(P \vee \neg P) \vdash \perp$ (de 17., $\perp$ -intro)
19. $\vdash (\neg(P \vee \neg P) \implies \perp)$ (de 18., $\implies$ -intro)
20. $\vdash \neg\neg(P \vee \neg P)$ (de 19., $\neg$ -intro)
21. $\vdash (P \vee \neg P)$ (de 20., $\neg\neg$ -elim).

On remarque ici que cette démonstration est très longue, lourde, et peu pratique. On peut mettre en place une idéographie pour rendre tout ça plus simple, mais le système de démonstration formelle est inhéremment long et "pénible".

Je conseille aux lectrices de s'exercer un peu en prouvant quelques formules valides comme $P \implies (Q \implies P)$, ou $P \implies \neg\neg P$, $(P \implies Q) \implies (\neg Q \implies \neg P)$ etc. Qu'il/elle sache que s'il/elle sait qu'une formule est valide, il/elle pourra la prouver. Rien ne dit en combien de temps, bien sûr, mais c'est toujours possible, précisément en vertu du théorème de complétude.

Avant d'attaquer la section suivante sur les algèbres de Boole qui permettra d'apporter une preuve du théorème, il nous faut vérifier quelque chose d'important: c'est ce que les anglophones appellent la *soundness* du système de preuve, i.e. qu'on ne peut prouver que des choses vraies. En effet, je peux facilement obtenir un théorème de complétude si j'ajoute comme règle qu'on peut inférer tout de n'importe quoi; le théorème ne sera alors pas intéressant.

Avant de s'attaquer à la complétude, il faut donc vérifier que si le séquent $\Gamma \vdash P$ est prouvable, alors $\Gamma \models P$. Autrement dit, si on peut prouver que P est vrai sous les hypothèses Γ , alors dès lors que les hypothèses sont vérifiées, P l'est aussi: un système de preuve qui ne vérifie pas ça serait

bien inutile et peu intéressant !

(Remarque: le théorème de complétude est la réciproque: si $\Gamma \models P$, alors $\Gamma \vdash P$. Si le système est bien construit, la preuve de sa *soundness* est aisée, mais la preuve de la complétude est en général plus difficile)

Théorème: Soit $\Gamma \cup \{P\} \subset \mathcal{F}$. Si $\Gamma \vdash P$ est démontrable, alors $\Gamma \models P$.

Preuve: La structure des démonstrations telles qu'on les a définies rend la tâche de démontrer ce théorème aisé. Avant la preuve, voici une idée de ce qu'elle va être: chacune des possibilités pour D_i dans la définition de démonstration peut être vue comme une règle d'inférence: de $\Gamma, P \vdash Q$ j'infère $\Gamma \vdash (P \implies Q)$.

Ainsi pour montrer qu'une démonstration ne peut former que des séquents "valides" (en un sens à définir plus tard) il suffit de montrer que chaque règle d'inférence conserve la validité. Au vu de la structure de la démonstration, on peut faire ça par récurrence sur une démonstration.

Un séquent $\Omega \vdash Q$ est dit valide si $\Omega \models Q$.

On suppose donc que D est une démonstration du séquent $\Gamma \vdash P$. On suppose que D est de longueur n . On montre par récurrence forte sur $i \leq n$ que D_i est valide, ce qui permettra de conclure avec $i = n$.

Soit donc $i \leq n$, on suppose que pour tout $j < i$ la propriété est vérifiée (à savoir "le séquent D_j est valide"). Alors par définition de démonstration, D_i est d'une des formes suivantes:

- $\Gamma \vdash P$ avec $P \in \Gamma$. Dans ce cas, par définition de $\models$, si δ est une valuation et $\delta \models \Gamma$, $\delta \models P$, donc $\Gamma \models P$.
- $\Gamma \vdash P$ et il existe $j < i$, Q tel que $\Gamma \vdash (P \wedge Q)$ (resp. $\Gamma \vdash (Q \wedge P)$) soit D_j . Puisque D_j est alors valide par hypothèse de récurrence, on en déduit que si δ est une valuation et si $\delta \models \Gamma$, alors $\delta \models (P \wedge Q)$, de sorte que $\bar{\delta}((P \wedge Q)) = 1$, et donc par définition, $\bar{\delta}(P) \geq \bar{\delta}((P \wedge Q)) = 1$, $\bar{\delta}(P) = 1$, et donc $\delta \models P$, de sorte que $\Gamma \models P$.
- $\Gamma \vdash (Q \vee R)$ et il existe $j < i$ tel que D_j soit $\Gamma \vdash Q$ (sans perte de généralité). Par hypothèse de récurrence, D_j est valide, de sorte que $\Gamma \models Q$. Ainsi si $\delta \models \Gamma$, $\bar{\delta}((Q \vee R)) \geq \bar{\delta}(Q) = 1$, donc $\delta \models (Q \vee R)$, de sorte que $\Gamma \models (P \vee Q)$.
- $\Gamma \vdash P$, et il existe $j, k, l < i$, Q, R tels que D_j soit $\Gamma \vdash (Q \vee R)$, D_k soit $\Gamma, Q \vdash P$ et D_l soit $\Gamma, R \vdash P$. Par hypothèse de récurrence, ces trois séquents sont valides. Soit alors δ une valuation telle que $\delta \models \Gamma$. On a alors $\delta \models (Q \vee R)$. Par une propriété établie précédemment, on en déduit que $\delta \models Q$ ou $\delta \models R$. Si $\delta \models Q$, alors $\delta \models \Gamma \cup \{Q\}$ par définition, et puisque $\Gamma, Q \vdash P$ par hypothèse (le séquent D_k étant valide), on en déduit que $\delta \models P$. Sinon, c'est que $\delta \models R$, on peut alors faire le même raisonnement et dans tous les cas obtenir $\delta \models P$, de sorte que finalement $\Gamma \models P$.

(il y a d'autres cas à examiner, mais on les laisse en exercice aux lecteurices, qui devrait pouvoir le faire sans difficulté au vu des cas traités. On fait simplement la remarque que si D_j est $\Gamma \vdash \perp$ avec $j < i$, puisque D_j est valide, toute valuation qui satisfait Γ satisfait aussi $\perp$: il n'y a donc pas de valuation qui satisfasse Γ , puisqu'il n'y a pas de valuation qui satisfasse $\perp$ par définition)

Avec l'exercice laissé aux lecteurices, cela conclut la preuve, on peut ainsi appliquer la propriété à $i = n$ et comme D_n est $\Gamma \vdash P$, on en déduit que ce séquent est valide: $\Gamma \models P$.

Notre système de démonstration n'est donc pas "trop fort": il ne prouve que des séquents valides (ouf!). Il ne reste plus qu'à voir qu'il est suffisamment fort, c'est-à-dire qu'il prouve tous les séquents valides. Pour cela il nous faudra passer par les algèbres de Boole, en particulier l'algèbre de Lindenbaum-Tarski d'une théorie T .

Avant de nous lancer là-dedans, voyons un peu comment interpréter nos séquents. Le séquent $\Gamma \vdash P$ peut se lire " Γ prouve P ". Si on observe la définition de démonstration, on observe que au cours d'une démonstration les hypothèses (Γ) changent. Essentiellement c'est à chaque fois une hypothèse qu'on fait puis qu'on *décharge*: je veux prouver $\neg P$, je rajoute donc P à mes hypothèses, j'en déduis une contradiction, et je peux donc conclure $\neg P$ en perdant l'hypothèse P . Si le séquent

final de la démonstration est $\Gamma \vdash P$, Γ est l'ensemble des hypothèses qu'on "voulait vraiment".

On appelle *théorie* tout sous-ensemble de $\mathcal{F}$. On dit qu'une formule P est conséquence syntaxique d'une théorie T , ce que l'on note à nouveau $T \vdash P$, si le séquent $T \vdash P$ est démontrable. Il est clair au vu de la définition de démonstration que $T \vdash P$ si et seulement si il existe un ensemble T_0 fini inclus dans T tel que $T_0 \vdash P$.

Notons alors $C_{\vdash} : \mathcal{P}(\mathcal{F}) \rightarrow \mathcal{P}(\mathcal{F})$ (resp. $C_{\models}$) l'opérateur $C_{\vdash}(T) = \{P \in \mathcal{F} \mid T \vdash P\}$ (resp. $C_{\models}(T) = \{P \in \mathcal{F} \mid T \models P\}$).

Ce que l'on vient de montrer c'est que pour tout T , $C_{\vdash}(T) \subset C_{\models}(T)$, ce qui est à espérer de tout système démonstratif. La partie intéressante, c'est la théorème de complétude qui revient à dire que l'inclusion est en fait une égalité.

Passons donc aux algèbres de Boole.

3 Les algèbres de Boole et l'algèbre de Lindenbaum-Tarski

3.1 Motivation

Un cadre naturel pour interpréter les formules propositionnelles est un ensemble muni d'opérations qui correspondent aux connecteurs propositionnels considérés lors de la définition de $\mathcal{F}$.

En effet si on a un ensemble X avec une opération binaire correspondant au $\wedge$, une correspondant au $\vee$, au $\implies$, au $\neg$, au $\top$ et $\perp$, alors pour tout choix de fonction $f : \mathcal{V} \rightarrow X$, on peut construire par récurrence comme on l'a fait pour $\{0, 1\}$ une extension à $\bar{f} : \mathcal{F} \rightarrow X$ qui vérifie toutes les égalités de la forme $\bar{f}((P \wedge Q)) = \bar{f}(P) \wedge \bar{f}(Q)$ etc.

Mais on veut aussi que ces opérations sur X se "comportent bien" de sorte que l'interprétation dans X ne soit pas complètement inutile. Toutes ces considérations mènent à la définition des algèbres de Boole.

3.2 Les algèbres de Boole

Définition: Une algèbre de Boole est un ensemble B muni d'opérations $\wedge, \vee : B \times B \rightarrow B$, $\neg : B \rightarrow B$, $\perp, \top \in B$ vérifiant des conditions détaillées plus tard. On note souvent $\mathfrak{B} := (B, \wedge, \vee, \neg, \top, \perp)$, et souvent par abus de notation on confond $\mathfrak{B}, B$, mais dans la plupart des cas ce n'est pas préjudiciable.

Les conditions à vérifier sont les suivantes: pour tous $x, y, z \in B$

- (distributivités) $x \wedge (y \vee z) = (x \wedge y) \vee (x \wedge z)$, $x \vee (y \wedge z) = (x \vee y) \wedge (x \vee z)$
- (associativités) $x \wedge (y \wedge z) = (x \wedge y) \wedge z$, $x \vee (y \vee z) = (x \vee y) \vee z$
- (commutativités) $x \wedge y = y \wedge x$, $x \vee y = y \vee x$
- (idempotences) $x \wedge x = x \vee x = x$
- (complémentaires) $x \wedge \neg x = \perp$, $x \vee \neg x = \top$
- (éléments neutres et absorbances) $x \wedge \top = x \vee \perp = x$, $x \vee \top = \top$, $x \wedge \perp = \perp$
- (lois de De Morgan) $\neg(x \wedge y) = \neg x \vee \neg y$, $\neg(x \vee y) = \neg x \wedge \neg y$
- (involutivité) $\neg \neg x = x$

Ces conditions sont souvent redondantes, mais on ne cherche pas ici à trouver de l'optimalité en termes de nombre de conditions. Elles peuvent sembler très abstraites, mais ces conditions se justifient très facilement si on considère au lieu de $=$ l'équivalence entre les formules propositionnelles: il est clair que le séquent $\vdash (((P \wedge P) \implies P) \wedge (P \implies (P \wedge P)))$ est prouvable par exemple, de sorte que $P \wedge P$ et P sont des formules "équivalentes" (de l'une on déduit l'autre et réciproquement). Cette remarque justifie par exemple les conditions d'idempotence. le.a lecteurice intrigué pourra montrer qu'en remplaçant $=$ par "sont équivalentes" dans les conditions précédentes et en interprétant les x, y, z comme des formules propositionnelles, toutes les conditions sont démontrables. Nous ne nous attarderons pas là-dessus ici.

Souvent on note $0, 1$ au lieu de $\perp, \top$ mais cela n'a pas d'importance. Un premier exemple (en fait

primordial) d'algèbre de Boole est $\{0, 1\}$. En effet si on le munit de $\min, \max, (x \mapsto 1 - x), 1, 0$, il devient une algèbre de Boole, et c'est assez facile à montrer (quoique légèrement long). On la note dans la suite $\mathfrak{B}_2$.

Une remarque qu'il est nécessaire de faire pour éviter les confusions: $\mathcal{F}$, muni de $(P, Q) \rightarrow (P \wedge Q)$ etc. n'est *pas* une algèbre de Boole. En effet, par exemple $(P \wedge P) \neq P$ (il suffit de regarder le nombre de caractères dans les deux formules pour s'en convaincre). Pourtant on va pouvoir le rapprocher des algèbres de Boole et utiliser ce rapprochement pour prouver notre théorème. Ce qui suit est un balbutiement de théorie des algèbres de Boole, nous développons juste les outils nécessaires pour la suite. A première vue, une bonne partie de ce qui suit n'a pas de lien avec notre théorème, mais pourtant, si.

Définition: Si $\mathfrak{B}, \mathfrak{C}$ sont deux algèbres de Boole, un morphisme d'algèbres de Boole $f : \mathfrak{B} \rightarrow \mathfrak{C}$ est une application $f : B \rightarrow C$ compatible avec les différentes opérations, c'est-à-dire $f(x \wedge y) = f(x) \wedge f(y), f(x \vee y) = f(x) \vee f(y), f(\neg x) = \neg f(x), f(\perp) = \perp, f(\top) = \top$, pour tous $x, y \in B$ (en commettant l'abus de notation consistant à noter de la même manière les opérations de $\mathfrak{B}$ et de $\mathfrak{C}$).

Comme dans beaucoup de cas en algèbre, on s'intéresse aux morphismes, et comme souvent on s'intéresse à certaines images réciproques. Ici, on s'intéresse à $f^{-1}(\{\perp\})$ pour un morphisme d'algèbres de Boole f .

Si $f(x) = f(y) = \perp$, alors $f(x \vee y) = f(x) \vee f(y) = \perp \vee \perp = \perp$. Si de plus $x \wedge y = x$ et $f(y) = \perp$, alors $f(x) = f(x \wedge y) = f(x) \wedge f(y) = f(x) \wedge \perp = \perp$. Finalement, si l'algèbre $\mathfrak{C}$ n'est pas "triviale", c'est-à-dire ne vérifie pas $\perp = \top$ (Exercice aux lecteurices: montrer que dans une algèbre de Boole, $\perp = \top$ si et seulement si l'algèbre de Boole n'a qu'un seul élément), alors $f(\top) = \top \neq \perp$, et donc $\top \notin f^{-1}(\{\perp\})$.

On va s'intéresser aux sous-ensembles de $\mathfrak{B}$ vérifiant ces conditions.

Définition: Soit $\mathfrak{B}$ une algèbre de Boole. Un sous-ensemble non vide $I \subset B$ est un *idéal* de $\mathfrak{B}$ si et seulement si il vérifie les conditions suivantes :

- Si $x, y \in I$ alors $x \vee y \in I$
- Si $x \wedge y = x$ et $y \in I$, alors $x \in I$
- $\top \notin I$

En fait, la notion d'idéal est conçue pour être équivalente à " I est un idéal si et seulement si il est de la forme $f^{-1}(\{\perp\})$ pour un certain morphisme d'algèbres de Boole f ". Par ce qui précède, on a déjà le "si".

Soit I un idéal, et soit $\mathcal{R}$ la relation définie sur B par $x \mathcal{R} y \iff \exists i \in I, x \vee i = y \vee i$. On montre dans la suite qu'il s'agit d'une relation dite de congruence, et on la notera $\equiv_I$.

Proposition: $\equiv_I$ est une relation d'équivalence. De plus, si $x \equiv_I x', y \equiv_I y'$ alors $x \wedge y \equiv_I x' \wedge y', x \vee y \equiv_I x' \vee y', \neg x \equiv_I \neg x'$.

Preuve: Puisque I est non vide (par définition), il existe $x \in I$. Puisque $x \wedge \perp = \perp$, par définition, $\perp \in I$.

Soit $x \in B$. Puisque $x \vee \perp = x \vee \perp$, $x \equiv_I x$. La relation $\equiv_I$ est clairement symétrique.

Finalement si $x \equiv_I y, y \equiv_I z$, alors on trouve $i, j \in I$ tels que $x \vee i = y \vee i, y \vee j = z \vee j$, et alors $x \vee (i \vee j) = (x \vee i) \vee j = (y \vee i) \vee j = (y \vee j) \vee i = (z \vee j) \vee i = z \vee (i \vee j)$. Or $i \vee j \in I$ car I est un idéal, donc $x \equiv_I z$. $\equiv_I$ est donc bien une relation d'équivalence.

Soit x, x', y, y' satisfaisant les hypothèses de l'énoncé, et $i, j \in I$ tels que $x \vee i = x' \vee i, y \vee j = y' \vee j$. Alors $x \vee y \vee (i \vee j) = x \vee i \vee y \vee j = x' \vee i \vee y' \vee j = x' \vee y' \vee (i \vee j)$, donc $x \vee y \equiv_I x' \vee y'$. On procède de même pour $x \wedge y$ en appliquant cette fois-ci la distributivité.

Finalement, $x \vee i = x' \vee i$ implique que $(x \vee i) \wedge \neg i = (x' \vee i) \wedge \neg i$, d'où par distributivité, par les règles de complémentaires, et par les règles d'éléments neutres, $x \wedge \neg i = x' \wedge \neg i$. En appliquant $\neg$ et en appliquant les lois de De Morgan et l'involutivité, on obtient $\neg x \vee i = \neg x' \vee i$, de sorte que $\neg x \equiv_I \neg x'$.

Cette proposition permet de définir une structure d'algèbre de Boole non triviale sur le quotient $B / \equiv_I$: notons $\pi : B \rightarrow B / \equiv_I$ la surjection canonique, qui à un élément de B associe sa classe d'équivalence dans la relation $\equiv_I$. On définit alors $\pi(x) \wedge \pi(y) := \pi(x \wedge y)$.

Ceci est bien défini car d'après ce qui précède, la classe d'équivalence de $x \wedge y$ ne dépend pas de x ni de y mais seulement de leurs classes d'équivalences respectives. De même on définit $\pi(x) \vee \pi(y) := \pi(x \vee y)$, $\neg\pi(x) := \pi(\neg x)$, $\perp := \pi(\perp)$, $\top := \pi(\top)$. Ces opérations sont bien définies pour la même raison. Comme de plus π est surjective, elles sont définies sur tout $B/\equiv_I$.

C'est un exercice routinier de vérifier à partir de là que les conditions d'algèbres de Boole sont vérifiées, car elles le sont dans B . L'algèbre obtenue est de plus non triviale car $\top \notin I$, donc $\top$ n'est pas équivalent à $\perp$, donc $\pi(\top) \neq \pi(\perp)$ (en effet si il existe $i \in I$ tel que $\perp \vee i = \top \vee i$, alors $i = \top$, donc $\top \in I$, or on a supposé le contraire).

La nouvelle algèbre de Boole obtenue est appelée algèbre quotient, et est notée $\mathfrak{B}/I$.

Il est aussi facile de vérifier qu'alors, par construction, π est un morphisme d'algèbres de Boole. Observons ce que vaut $\pi^{-1}(\{\perp\})$: comme le $\perp$ de $\mathfrak{B}/I$ est défini par $\pi(\perp)$, si $\pi(x) = \perp$, c'est que $\pi(x) = \pi(\perp)$, soit $x \equiv_I \perp$. Mais si $x \equiv_I \perp$, il existe $i \in I$ tel que $x \vee i = \perp$. Donc $x \wedge i = x \wedge (x \vee i) = x$ (exercice laissé aux lectrices : $x \wedge (x \vee y) = x$), de sorte que puisque I est un idéal, $x \in I$. Réciproquement, si $x \in I$, $x \equiv_I \perp$ donc $\pi(x) = \perp$. Ainsi $I = \pi^{-1}(\{\perp\})$.

Ainsi on a :

Proposition: Un sous-ensemble $I \subset B$ d'une algèbre de Boole $\mathfrak{B}$ est un idéal si et seulement s'il existe un morphisme d'algèbres de Boole $f : \mathfrak{B} \rightarrow \mathfrak{C}$ tel que $I = f^{-1}(\{\perp\})$.

En étant arrivés là, je peux décrire le chemin que nous emprunteront pour démontrer le théorème de complétude.

On considère une théorie T , à laquelle on va associer une algèbre de Boole $\mathfrak{B}_T$ (dite algèbre de Lindenbaum-Tarski) qui reflètera la démontrabilité dans T . On montrera alors que si une proposition P n'est pas prouvable, elle est contenue dans un idéal spécial I de $\mathfrak{B}_T$ tel que $\mathfrak{B}_T/I$ est essentiellement $\mathfrak{B}_2$. De ceci on déduira une valuation qui envoie P sur 0, ce qui prouvera que P n'est pas conséquence sémantique de T . Il nous reste donc à voir les algèbres de Lindenbaum-Tarski, les idéaux dits premiers ou maximaux (dans le cas des algèbres de Boole, ces deux notions coïncident), et nous verrons en quoi cela fournit une valuation appropriée.

Soit I un idéal de $\mathfrak{B}$. On dit qu'il est *maximal* s'il n'est inclus strictement dans aucun autre idéal de $\mathfrak{B}$.

Proposition : Un idéal est maximal si et seulement si pour tout $x \in B$, $x \in I$ ou $\neg x \in I$.

Preuve: Si la condition est vérifiée, supposons que $I \subset J$, $I \neq J$ où J est aussi un idéal. Soit $x \in J \setminus I$. Comme $x \notin I$, $\neg x \in I$. Donc $\neg x, x \in J$. Or J est un idéal donc $x \vee \neg x = \top \in J$: contradiction.

Réciproquement, supposons que $x, \neg x \notin I$. Supposons qu'il existe $i \in I$ tel que $i \vee x = \top$. Alors $(i \vee x) \wedge \neg x = \neg x$ et donc $i \wedge \neg x = \neg x$, et comme I est un idéal, $\neg x \in I$, ce qui contredit l'hypothèse, donc pour tout $i \in I$, $i \vee x \neq \top$. De même, pour $i \in I$, $i \vee \neg x \neq \top$. Considérons alors par $J = \{y \in B \mid \exists i \in I, y \wedge (i \vee x) = y\}$.

Comme $i \wedge (i \vee x) = i$ pour tout $i \in I$, on a $I \subset J$. De même, $x \in J$ donc $I \neq J$. De plus, d'après ce qui précède, $\top \notin J$.

Supposons $y, z \in J$, et soit $i, j \in I$ tels que $y \wedge (i \vee x) = y$ et $z \wedge (j \vee x) = z$. Calculons $(y \vee z) \wedge (i \vee j \vee x)$:

$(y \vee z) \wedge (i \vee j \vee x) = ((i \vee j \vee x) \wedge y) \vee ((i \vee j \vee x) \wedge z) = (((i \vee x) \wedge y) \vee (j \wedge y)) \vee (((j \vee x) \wedge z) \vee (i \wedge z)) = (y \vee (j \wedge y)) \vee (z \vee (i \wedge z)) = y \vee z$. Donc $y \vee z \in J$.

Supposons $y \wedge t = t$. Soit i tel que $y \wedge (i \vee x) = y$. Alors $t \wedge (i \vee x) = t \wedge y \wedge (i \vee x) = t \wedge y = t$, donc $t \in J$.

On a donc montré que J était un idéal qui contenait strictement I : I n'est pas maximal. Donc si I est maximal, pour tout x , $x \in I$ ou $\neg x \in I$.

Si un morphisme d'algèbres de Boole est bijectif, on appelle ça un isomorphisme. Soit $f : \mathfrak{B} \rightarrow \mathfrak{C}$ un tel morphisme, et soit f^{-1} sa réciproque. Alors on vérifie facilement que f^{-1} est aussi un morphisme d'algèbres de Boole. S'il existe un isomorphisme entre $\mathfrak{B}$ et $\mathfrak{C}$ on dit que ces deux algèbres sont isomorphes.

Deux algèbres isomorphes sont identiques pour tout ce qui est intéressant en termes d'algèbres de Boole.

Proposition: Soit $\mathfrak{B}$ une algèbre de Boole et I un idéal maximal de $\mathfrak{B}$. Alors $\mathfrak{B}/I$ est isomorphe à $\mathfrak{B}_2$.

Preuve : On va montrer qu'une algèbre de Boole à deux éléments est isomorphe à $\mathfrak{B}_2$, et que sous les hypothèses de l'énoncé, $\mathfrak{B}/I$ a deux éléments, ce qui permettra de conclure.

Soit donc $\mathfrak{B}$ une algèbre de Boole à deux éléments. D'après un exercice qu'on avait laissé aux lectrices, $\top \neq \perp$ dans $\mathfrak{B}$. Donc l'un des deux éléments est $\top$, l'autre est $\perp$. Donc toutes les opérations $\wedge, \vee, \neg, \perp$ sont entièrement déterminées. De plus, $\neg\top = \perp, \neg\perp = \top$ (on peut le montrer facilement dans une algèbre de Boole quelconque, mais ici c'est encore plus immédiat: la lectrice vérifiera ce qu'il/elle souhaite). Donc l'application $\mathfrak{B} \rightarrow \mathfrak{B}_2$ définie par $\top \mapsto 1, \perp \mapsto 0$ est une morphisme d'algèbres de Boole qui est clairement bijectif: c'est donc un isomorphisme.

Soit $\mathfrak{B}, I$ tels que dans l'énoncé. Soit $x \in B \setminus I$. Puisque $x \notin I$, et car I est maximal, d'après la proposition précédente, $\neg x \in I$. Donc $x \vee \neg x = \top = \top \vee \neg x$, de sorte que $x \equiv_I \top$. Si $x \in I$ clairement $x \equiv_I \perp$. Donc il y a précisément deux classes d'équivalences dans $B/\equiv_I$, celle de $\perp$ et celle de $\top$, donc $\mathfrak{B}/I$ est une algèbre de Boole à deux éléments, donc elle est isomorphe (d'après ce qui précède) à $\mathfrak{B}_2$.

3.3 L'algèbre de Lindenbaum-Tarski de T et le théorème de complétude

Si on a une application $g : \mathcal{F} \rightarrow \{0, 1\}$ qui respecte les opérations $\wedge, \vee, \neg, \implies, \top, \perp$ (où $\implies$ est défini dans une algèbre de Boole comme $x \implies y := \neg x \vee y$, donc le respect de $\implies$ est conséquence de celui de $\neg$ et $\vee$), elle est nécessairement de la forme $g = \bar{f}$, où $f : \mathcal{V} \rightarrow \{0, 1\}$ est une valuation quelconque.

Ainsi si on a une théorie T , et une formule propositionnelle P telle que le séquent $T \vdash P$ n'est pas démontrable, il nous suffit de trouver une telle application $g : \mathcal{F} \rightarrow \{0, 1\}$ telle que $g(P) = 0$, et $g(Q) = 1$ pour tout $Q \in T$ pour conclure que P n'est pas conséquence sémantique de T . On va construire cette application en composant plusieurs: on va d'abord composer par une surjection canonique vers une algèbre de Boole (celle dont je parle depuis un moment: celle de Lindenbaum-Tarski), puis par une surjection canonique de celle-ci vers un quotient par un idéal maximal, puis par un isomorphisme avec $\mathfrak{B}_2$. Chacune de ces applications préservera les différentes opérations considérées, de sorte que la composition les préservera aussi. La construction des applications sera telle que $g(P) = 0$ et $g(T) \subset \{1\}$.

Définition: Soit T une théorie. On définit alors la relation suivante sur $\mathcal{F}$: $P \equiv_T Q$ si et seulement si le séquent $T \vdash ((P \implies Q) \wedge (Q \implies P))$ est démontrable.

Proposition: Pour toute théorie T , $\equiv_T$ est une relation d'équivalence. Si de plus $P \equiv_T P', Q \equiv_T Q'$, alors $(P \wedge Q) \equiv_T (P' \wedge Q'), (P \vee Q) \equiv_T (P' \vee Q'), (P \implies Q) \equiv_T (P' \implies Q')$, et $\neg P \equiv_T \neg P'$. De plus, $T \vdash P$, si et seulement si $P \equiv_T \top$.

Preuve: Remarquons d'abord (c'est clair d'après la définition de démonstration) que $P \equiv_T Q$ si et seulement si les deux séquents $T, P \vdash Q$ et $T, Q \vdash P$ sont démontrables. Cette remarque simplifiera la suite.

Il est ainsi clair que $\equiv_T$ est symétrique (la deuxième caractérisation l'est) et réflexive (d'après la deuxième caractérisation et la règle Ax.). Supposons $P \equiv_T Q, Q \equiv_T R$. De $T, Q \vdash R$ on infère $T \vdash (Q \implies R)$, et donc clairement $T, P \vdash (Q \implies R)$ est démontrable. Or $T, P \vdash Q$ est démontrable par hypothèse, donc par MP., on peut inférer $T, P \vdash R$. On montre de même que $T, R \vdash P$, de sorte que $P \equiv_T R$.

On laisse aux lectrices le soin de prouver les autres clauses de la proposition, en utilisant la deuxième caractérisation de $\equiv_T$, et en appliquant la définition de démonstration.

On laisse aussi le soin aux lectrices de prouver que les séquents sont démontrables :

- $\vdash ((P \wedge (Q \vee R)) \implies ((P \wedge Q) \vee (P \wedge R)))$ et $\vdash (((P \wedge Q) \vee (P \wedge R)) \implies (P \wedge (Q \vee R)))$
- $\vdash ((P \vee (Q \wedge R)) \implies ((P \vee Q) \wedge (P \vee R)))$ et $\vdash (((P \vee Q) \wedge (P \vee R)) \implies (P \vee (Q \wedge R)))$
- $\vdash ((P \wedge (Q \wedge R)) \implies ((P \wedge Q) \wedge R))$ et $\vdash (((P \wedge Q) \wedge R) \implies (P \wedge (Q \wedge R)))$
- $\vdash ((P \vee (Q \vee R)) \implies ((P \vee Q) \vee R))$ et $\vdash (((P \vee Q) \vee R) \implies (P \vee (Q \vee R)))$
- $P \wedge Q \vdash Q \wedge P$ et $P \vee Q \vdash Q \vee P$

- $P \wedge P \vdash P, P \vdash P \wedge P, P \vee P \vdash P, P \vdash P \vee P$
- $P \wedge \neg P \vdash \perp, \perp \vdash P \wedge \neg P, P \vee \neg P \vdash \top, \top \vdash P \vee \neg P$
- $P \wedge \top \vdash P, P \vdash P \wedge \top, P \vee \perp \vdash P, P \vdash P \vee \perp$
- $P \vee \top \vdash \top, \top \vdash P \vee \top, P \wedge \perp \vdash \perp, \perp \vdash P \wedge \perp$
- $\neg\neg P \vdash P, P \vdash \neg\neg P$
- $\neg(P \wedge Q) \vdash (\neg P \vee \neg Q), (\neg P \vee \neg Q) \vdash \neg(P \wedge Q)$
- $\neg(P \vee Q) \vdash (\neg P \wedge \neg Q), (\neg P \wedge \neg Q) \vdash \neg(P \vee Q)$

et d'en déduire que l'ensemble quotient $\mathcal{F}/\equiv_T$, muni de $\wedge, \vee, \neg, \top, \perp$ (dont il faut vérifier qu'elles sont bien définies) est une algèbre de Boole.

le.a lecteurice attentif/ve aura remarqué que démontrer tous ces séquents revient essentiellement à montrer cela (vous aurez remarqué l'analogie entre ces séquents et les conditions qui définissent une algèbre de Boole).

L'algèbre de Boole obtenue, que l'on note $\mathfrak{B}_T$ est l'algèbre de Lindenbaum-Tarski de T .

D'après la proposition qui précède sa définition, on voit que, notant $\pi : \mathcal{F} \rightarrow \mathfrak{B}_T$ la surjection canonique, $T \vdash P$ si et seulement si $\pi(P) = \top$. Ainsi cette algèbre donne une caractérisation algébrique de la démontrabilité par T .

Il est clair que π conserve les opérations $\wedge, \vee, \neg, \perp, \top$, de sorte qu'on a fait le premier pas dans la construction de notre application $\mathcal{F} \rightarrow \{0, 1\}$.

Le reste de la construction dépend de la formule P considérée, et un peu de T .

On suppose que le séquent $T \vdash \perp$ n'est pas démontrable, ce qui revient à dire que dans $\mathfrak{B}_T$, $\perp \neq \top$. Si ce n'est pas le cas, c'est que $\mathfrak{B}_T$ n'a qu'un élément, donc de toute façon $T \vdash P$: la conclusion du théorème de complétude est dans ce cas vérifiée, donc on n'a rien à dire.

Donc on peut supposer que $T \vdash \perp$ n'est pas démontrable, soit donc que $\top \neq \perp$ dans $\mathfrak{B}_T$, donc que $\mathfrak{B}_T$ a plusieurs éléments.

Soit $P \in \mathcal{F}$ tel que $T \vdash P$ ne soit pas démontrable. Ainsi, $\pi(P) \neq \top$.

C'est à cet instant précis que l'on utilise le résultat mentionné dans l'introduction: le *Boolean Prime Ideal Theorem*, en court BPI. Nous utilisons ici un cas particulier de principe, qu'on prend comme axiome. En appendice, on pourra voir qu'en supposant le théorème de complétude, on peut le démontrer, de sorte qu'avec ZF en théorie ambiante, ces deux principes sont équivalents. Une "justification" du BPI est qu'il est conséquence de l'axiome du choix, usuellement accepté dans l'axiomatisation. Nous énonçons ici le cas particulier dont on a besoin:

Axiome (théorème de ZFC): Pour toute algèbre de Boole non réduite à un élément $\mathfrak{B}$ et tout $a \in B \setminus \{\top\}$, il existe un idéal maximal de $\mathfrak{B}$ contenant a .

Nous avons désormais tous les outils en main pour conclure: nous savons que $\pi(P) \neq \perp$ dans $\mathfrak{B}_T$, considérons donc un idéal maximal I de $\mathfrak{B}_T$ contenant $\pi(P)$. On a alors un morphisme surjectif canonique $\zeta : \mathfrak{B}_T \rightarrow \mathfrak{B}_T/I$. $\zeta(\pi(P)) = \perp$ car $\pi(P) \in I$ par définition.

Soit alors $\xi : \mathfrak{B}_T/I \rightarrow \mathfrak{B}_2$ un isomorphisme (d'après une proposition précédente, il en existe un), de sorte que $\xi(\zeta(\pi(P))) = \xi(\perp) = 0$.

On a alors une application $g := \xi \circ \zeta \circ \pi : \mathcal{F} \rightarrow \{0, 1\}$ qui est une composée de deux morphismes d'algèbres de Boole et d'une application qui préserve $\wedge, \vee, \neg, \perp, \top$ (et donc $\implies$), qui préserve donc aussi ces opérations, qui est donc de la forme $g = \bar{f}$ pour une certaine valuation f ; telle que $g(P) = 0$. Donc $f \not\models P$.

Il ne reste plus qu'à vérifier que $f \models T$. Soit donc $Q \in T$. En particulier, $T \vdash Q$ (par Ax.), de sorte que par les propriétés établies précédemment, $Q \equiv_T \top$, donc $\pi(Q) = \top$, et comme ζ et ξ sont des morphismes d'algèbres de Boole, $g(Q) = \xi \circ \zeta \circ \pi(Q) = \top$. Donc $f \models T, f \not\models P$: $T \not\models P$.

On a donc montré:

Théorème de Complétude : Soit T une théorie et P une formule propositionnelle. Si $T \models P$, alors $\overline{T} \vdash P$.

Dans ce qui suit nous donnerons d'abord plusieurs reformulations immédiates mais intéressantes de ce théorème, puis nous donnerons comme application le fameux théorème de compacité.

Théorème de Complétude (version 1bis): Soit T une théorie et P une formule propositionnelle. Alors $T \models P$ si et seulement si $T \vdash P$.

On dit qu'une théorie T est consistante s'il existe une valuation δ telle que $\delta \models T$. On dit qu'elle est cohérente si le séquent $T \vdash \perp$ n'est pas démontrable (équivalamment "s'il existe un séquent de la forme $T \vdash P$ qui n'est pas démontrable": laissé en exercice). Une autre reformulation, frappante, est alors :

Théorème de Complétude (version 2): Une théorie est consistante si et seulement si elle est cohérente.

Preuve: Le fait que la consistance implique la cohérence vient du fait que tout séquent démontrable est valide, donc si $\delta \models T$, puisque $\delta \not\models \perp$, on ne peut avoir $T \vdash \perp$.

La réciproque vient de la considération suivante: soit T une théorie non consistante. Alors $T \models \perp$: toute valuation qui satisfait T satisfait $\perp$ (il n'y en a pas !) . Par le théorème de complétude (version 1bis), $T \vdash \perp$: T n'est pas cohérente.

Pour obtenir la version 1bis du théorème de complétude à partir de cette version-ci on remarque que la chose suivante: soit T une théorie, P une formule propositionnelle. Si T est incohérente, elle est inconsistante (par la version 2), donc $T \models P$, donc "si $T \vdash P$, $T \models P$ " est trivialement vrai. Si T est cohérente, et $T \vdash P$, alors $T \cup \{\neg P\}$ est incohérente, donc elle est inconsistante (par la version 2), donc toute valuation satisfaisant T ne satisfait pas $\neg P$: elle satisfait P ; donc $T \models P$. Donc dans tous les cas, on a le résultat attendu.

Réciproquement, supposons $T \not\models P$. Alors $T \cup \{\neg P\}$ est cohérente. Donc elle est consistante: une valuation qui la satisfait satisfait T et $\neg P$ et ne satisfait donc pas P : donc $T \not\models P$.

Les deux versions du théorème sont donc équivalentes.

Maintenant, une application directe du théorème de complétude :

Théorème de Compacité: Soit T une théorie inconsistante. Alors il existe $T_0 \subset T$ finie qui est inconsistante. Equivalamment, si une théorie a toutes ses sous-théories finies consistantes, alors elle l'est.

Preuve: Si $T \vdash \perp$, alors soit D une démonstration de ce séquent. Cette démonstration est une suite finie : soit T_0 la théorie constituée des différentes formules $P \in T$ tel qu'un séquent de la forme $\Gamma \vdash P$ apparaît dans D . Alors $T_0 \subset T$, et le séquent $T_0 \vdash \perp$ est prouvable, de sorte que T_0 est incohérente donc inconsistante.

4 Appendices

On propose dans la suite 3 appendices qui apportent un éclairage différent sur certains points de ce qui précède et aideront les lecteurs les plus avancés à mieux comprendre le sujet. Contrairement au reste du texte, ces appendices ne sont pas *self-contained* et font référence à des mathématiques non nécessairement développées ici. Ils seront donc moins accessibles que le reste du texte. En principe ils ne sont cependant bien évidemment pas nécessaire à la compréhension du corps principal du texte.

4.1 Algèbre Universelle

Si le lecteur connaît un peu d'algèbre universelle, une bonne partie de ce qui est traité ici peut être vu comme un cas particulier de constructions générales, ce qui simplifie ici quelques arguments, que l'on n'a ainsi pas à reproduire.

En effet, $\mathcal{F}$ peut être définie comme l'algèbre de type $\tau = \langle 2, 2, 2, 1, 0, 0 \rangle$ absolument libre sur N_0 générateurs (si on se restreint au cas $\mathcal{V}$ dénombrable : sinon c'est l'algèbre de type τ absolument libre sur $\mathcal{V}$ générateurs). Ceci facilite la définition de l'extension d'une valuation, mais aussi quelques arguments à ce sujet.

De plus on peut alors raccourcir l'expression "préserve telle et telle opération" en simplement "est un τ -morphisme".

Avec des bases en algèbre universelle, la construction du quotient d'une algèbre de Boole par un idéal n'est de plus plus mystérieuse, et toute cette partie peut disparaître.

Le fait qu'une algèbre de Boole quotientée par un idéal maximal donne une algèbre isomorphe à $\mathfrak{B}_2$ provient du fait que dans la variété des algèbres de Boole, congruences et idéaux correspondent bijectivement (aux congruences triviales près), et donc un quotient par un idéal maximal n'a que deux congruences: il est donc subdirectement irréductible, et il est connu que la seule algèbre de Boole subdirectement irréductible est $\mathfrak{B}_2$. L'existence du morphisme $\mathcal{F} \rightarrow \{0, 1\}$ envoyant P sur 0 n'est alors plus qu'une formalité, et le théorème devient évident: il suffit de vérifier que $\equiv_T$ est une congruence pour toute théorie T , mais ceci provient facilement du système de démonstration mis en place. En fait on pourrait même remarquer, bien plus simplement que si T est une théorie, alors $\pi(T)$ est un idéal de $\mathcal{F}/\equiv$ (où $\equiv := \equiv_T$ avec $T = \emptyset$), ce qui est encore plus simple (la vie est plus simple une fois qu'on a quotienté par l'équivalence logique...).

Ainsi une connaissance d'algèbre universelle permet d'aller plus vite dans cette preuve; et on perçoit mieux comment la généraliser à d'autres logiques (j'en parlerai rapidement dans l'appendice sur les autres logiques).

4.2 D'autres logiques

Il y a différentes manières d'aborder un point de vue plus général sur ce genre de question. La première généralisation évidente est de supprimer des règles d'inférence, ce qui "affaiblit" la logique. Les règles que j'ai données pour la logique classique sont largement sous-optimales, donc ici on pourrait en enlever et se retrouver avec la même logique.

En enlevant assez on se retrouve par exemple avec la logique intuitionniste, où $\mathcal{F}/\equiv$ n'est plus une algèbre de Boole, mais c'est une algèbre de Heyting (et à ce moment le $\implies$ devient -en général- différent du $\neg - \vee -$).

On peut aussi modifier le langage, par exemple avoir différentes disjonctions et conjonctions comme en logique linéaire, etc.

On peut ainsi obtenir une grande variété de logiques, qui sont étudiées notamment dans le cadre de la logique algébrique abstraite. On étudie alors la complétude d'une logique vis-à-vis d'une classe d'algèbres.

Dans le cas de la logique classique, cette classe est celle des algèbres de Boole, il s'avère que (par chance ?- la variété **Bool** est générée par une seule algèbre) la complétude par rapport à cette classe implique la complétude par rapport à $\mathfrak{B}_2$. Au contraire, la logique intuitionniste, elle est complète par rapport à la classe des algèbres de Heyting, mais pas par rapport à $\mathfrak{B}_2$ (sinon on pourrait conclure que les théorèmes classiques sont aussi théorèmes intuitionnistes, et on sait que c'est faux).

L'étude de ces logiques peut avoir différents intérêts, que cette étude n'ait pas de motivation ou parce que ces logiques peuvent être la logique interne de certaines structures qui intéressent les mathématiciens par d'autres aspects. Par exemple, la logique interne d'un topos n'est en général pas booléenne, mais elle est intuitionniste. On peut trouver des interprétations similaires à la logique linéaire par exemple (souvent, on s'intéresse à des logiques qui sont la logique interne d'une classe de catégories).

4.3 Interprétation topologique du théorème de compacité

Le dernier théorème présenté s'appelle le théorème de compacité. On pourrait se dire que ce nom est juste une analogie avec les espaces topologiques, l'analogie étant qu'on peut passer de quelque chose d'arbitraire (en particulier, potentiellement infini) à quelque chose de fini. On pourrait s'arrêter là. Mais il s'avère que ce théorème a une véritable interprétation topologique, qui dépasse la "simple" analogie, c'est-à-dire qu'on peut le démontrer indépendamment du théorème de complétude, en considérant certains espaces compacts.

En effet, $\{0, 1\}$ muni de la topologie discrète est un espace compact (au sens francophone du terme, c'est-à-dire que c'est un espace T_2 qui vérifie la propriété de Borel-Lebesgue), donc toute puissance de cet espace est aussi compacte: c'est le théorème de Tychonov (remarque: lui aussi requiert dans sa preuve le BPI. Si on veut le rendre encore plus fort et le prouver pour les espaces quasi-compacts, il faut utiliser l'axiome du choix). Ainsi $\{0, 1\}^\mathcal{V}$ est compact.

Soit T une théorie. On suppose que toute théorie finie $T_0 \subset T$ est consistante. Pour T_0 finie, le sous-ensemble $S(T_0) := \{\delta \in \{0, 1\}^\mathcal{V} \mid \delta \models T_0\}$ est un fermé (on peut l'exprimer comme une union

de produits de $\{0, 1\}$ avec un ensemble fini de $\{0\}$ et de $\{1\}$, car seul un nombre fini de variables apparaissent dans T_0 .

L'ensemble $\{S(T_0) \mid T_0 \subset T, T_0 \text{ finie}\}$ est stable par intersections finies car $\bigcap_{i=1}^n S(T_i) = S(\bigcup_{i=1}^n T_i)$,

et ne contient pas l'ensemble vide par hypothèse, et il ne contient que des fermés. Par compacité, son intersection est non vide. Or, clairement, son intersection est l'ensemble des valuations qui satisfont T . Donc T est consistante: on a obtenu ainsi le théorème de compacité.

4.4 Le BPI

L'énoncé précis et le plus fort du BPI est le suivant :

BPI: Soit $\mathfrak{B}$ une algèbre de Boole, I et un idéal et F un filtre disjoint de I . Alors I est inclus dans un idéal maximal disjoint de F .

Il est (peut-être étonnamment) équivalent au fait que toute algèbre de Boole non triviale admette un idéal maximal (on quotiente par les bonnes relations, puis on remonte à $\mathfrak{B}$ un idéal maximal du quotient). Il est alors facile de le déduire du lemme de Zorn, et donc en particulier de l'axiome du choix.

On peut montrer cependant que l'axiome du choix n'est pas conséquence du BPI (sous ZF).

Le BPI est équivalent au théorème de Tychonov (concernant les espaces compacts) et au théorème de compacité, ainsi qu'au théorème de complétude. L'axiome du choix est strictement plus fort et est, lui, équivalent au théorème de Tychonov concernant les espaces quasi-compacts.

Nous avons déjà vu que dans ZF, le BPI implique le théorème de complétude, qui implique le théorème de compacité. Il ne reste plus qu'à montrer que le théorème de compacité implique le BPI. C'est ce que nous ferons dans cet appendice.

On va donc montrer que toute algèbre de Boole non triviale admet un idéal maximal. Soit donc $\mathfrak{B}$ une telle algèbre de Boole. On considère un ensemble de variables propositionnelles $\mathcal{V}$, une variable P_a pour tout $a \in B$. Essentiellement on cherche une valuation δ telle que si on décide que $a \in I \iff \delta(P_a) = 1$, alors I sera un idéal maximal.

Soit T la théorie suivante: ses éléments sont ceux de la forme $(P_a \wedge P_b) \implies P_{a \vee b}$, pour $a, b \in B$, $P_a \implies P_b$ pour $a, b \in B$, si $a \wedge b = b$, $\neg P_\top$, $P_a \vee P_{\neg a}$ pour $a \in B$.

Il est alors clair que si $\delta \models T$, et si on définit $I := \{a \in B \mid \delta \models P_a\}$, alors I est un idéal maximal. Il suffit donc de montrer que T est consistante pour conclure. Par le théorème de compacité, il suffit de montrer que toute sous-théorie finie l'est.

Soit T_0 une théorie finie, et $K = \{a \in B \mid P_a \text{ apparaît dans } T_0\}$. K est un ensemble fini. L'algèbre de Boole qu'il génère est donc finie (car les algèbres de Boole libres sur un nombre fini de générateurs sont finies), on la note $\mathfrak{K}$. Elle admet un idéal $(\{\perp\})$ et donc clairement elle en admet un maximal I (on prend un idéal de cardinal maximal, il est alors maximal). En définissant $\delta(P_a) = 1$ si $a \in I$, 0 sinon, on obtient une valuation qui satisfait T_0 . Ainsi, T_0 est consistante.

Donc T est consistante, et $\mathfrak{B}$ a un idéal maximal: on a ainsi prouvé le BPI, et la boucle est bouclée.