



What Does This Notation Mean Anyway?

David A Feller, Joe B. B Wells, Fairouz Kamareddine, Sebastien Carlier

► To cite this version:

David A Feller, Joe B. B Wells, Fairouz Kamareddine, Sebastien Carlier. What Does This Notation Mean Anyway?: MBNF-Style Notation as it is Actually Used. 13th international Workshop on Logical Frameworks and Meta-Languages: Theory and Practice, Jul 2018, Oxford, United Kingdom. hal-01812800

HAL Id: hal-01812800

<https://inria.hal.science/hal-01812800>

Submitted on 11 Jun 2018

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution - NoDerivatives 4.0 International License

What Does This Notation Mean Anyway?

BNF-Style Notation as it is Actually Used

D. A. Feller

J. B. Wells

Sébastien Carlier

F. Kamareddine

Following the introduction of BNF notation by Backus for the Algol 60 report and subsequent notational variants, a metalanguage involving formal “grammars” has developed for discussing structured objects in Computer Science and Mathematical Logic. We refer to this offspring of BNF as *Math-BNF* or *MBNF*, to the original BNF and its notational variants just as *BNF*, and to aspects common to both as *BNF-style*. What all BNF-style notations share is the use of production rules roughly of this form:

$$\bullet ::= \circ_1 \mid \cdots \mid \circ_n$$

Normally, such a rule says “every instance of $\circ_i$ for $i \in \{1, \dots, n\}$ is also an instance of $\bullet$ ”.

MBNF is distinct from BNF in the entities and operations it allows. Instead of strings, MBNF builds arrangements of symbols that we call *math-text*. Sometimes “syntax” is defined by interleaving MBNF production rules and other mathematical definitions that can contain chunks of math-text.

There is no clear definition of MBNF. Readers do not have a document which tells them how MBNF is to be read and must learn MBNF through a process of cultural initiation. To the extent that MBNF is defined, it is largely through examples scattered throughout the literature.

This paper gives MBNF examples illustrating some of the differences between MBNF and BNF. We propose a definition of *syntactic math text* (SMT) which handles many (but far from all) uses of math-text and MBNF in the wild. We aim to balance the goal of being accessible and not requiring too much prerequisite knowledge with the conflicting goal of providing a rich mathematical structure that already supports many uses and has possibilities to be extended to support more challenging cases.

1 Background and Motivation

Understanding MBNF is important to interpreting papers in theoretical computer science. Out of the 30 papers in the ESOP 2012 proceedings [18], 19 used MBNF, while not one used BNF.¹ This section highlights some of the ways in which the notation we call MBNF differs from BNF. This should demonstrate that a definition could be helpful.

Where BNF uses Strings, MBNF Uses Math-Text In addition to arranging symbols from left to right on the page, math-text allows subscripting, superscripting, and placing text above or below other text. It also allows for marking whole segments of text, for example with an overbar (a vinculum). Readers can find more detailed information on how math-text can be laid out in The TeXbook [11], or the Presentation MathML [8] and OpenDocument [9] standards. Here is a nonsense piece of Math-text to illustrate how it may be laid out:

$${}^c\downarrow a' = \check{p}\langle v''_x \odot a^{2+1} \rangle - \overline{f_x^n + y \cdot f j} + \sum_{i=0}^{\infty} s_{i \in 1 \dots n} \xrightarrow{a, b, c} b \hat{a}$$

¹We chose ESOP 2012 because its book was the most recent conference proceedings that we had as a paper book. Because the first book we picked contained an abundance of challenging instances of MBNF, our wider searching has mainly been to find even more challenging examples. We will be happy to receive pointers to additional interesting cases.

Instead of non-terminal symbols, MBNF uses *metavariables*², which appear in math-text and obey the conventions of mathematical variables. Metavariables are not distinguished from other symbols by annotating them as BNF does, but by font, spacing, or merely tradition.

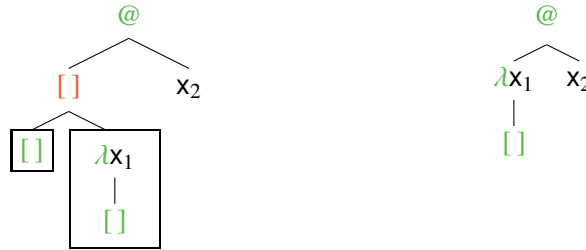
Parentheses for disambiguation are not needed in MBNF grammars and when an MBNF grammar specifies such parentheses they can often be omitted without any need to explain. When possible, MBNF takes advantage of the tree-like structure implicit in the layout of symbols on the page when features like superscripting and overbarring are used.

MBNF Is Aimed at Human Readers MBNF is meant to be interpreted by humans, not computers/parser generators. It is common to define a MBNF grammar in an article for humans and a separate EBNF grammar for use with a parser generator to build a corresponding implementation. Entities defined with MBNF are not intended or expected to be serialized or parsed and MBNF grammars are typically missing features needed to disambiguate complex terms. Papers often put complicated uses of the mathematical metalanguage in the middle of MBNF notation.

MBNF Allows Powerful Operators Like Context Hole Filling (a.k.a. Tree Splicing) Chang and Felleisen [2, p 134] present an MBNF grammar defining the λ -term contexts with one hole where the spine³ is a balanced segment⁴ ending in a hole. For explanatory purposes we alter their grammar slightly by writing $e@e$ instead of ee and adding parentheses. Concrete syntax and BNF-style notation are green. Metavariables are blue. Additional operators are red.

$$\begin{aligned} e &::= x \mid (\lambda x.e) \mid (e@e) \\ A &::= [] \mid (A[(\lambda x.A)]@e) \end{aligned}$$

One can think of the context hole filling operation in this grammar ($[]$ in $(A[(\lambda x.A)]@e)$) as performing tree splicing operations within the syntax. Consider these trees which illustrate steps in building syntax trees for A :



These trees show the result of the second rule where each A is $[]$ and e is a variable. The tree on the left is the tree corresponding to $A[\lambda x.A]@e$ before the hole filling operation is performed, where the first A is assigned $[]$. The tree on the right represents an unparsing of what we would normally consider the syntax tree for $((\lambda x_1.[])@x_2)$. We write x_1 and x_2 for disambiguated instances of x . A metavariable assigned a value won't appear in the final tree. If it's not a terminal node, $[]$ tells us to fill in the leaf in the frame on the left with the tree in the frame on the right. Once performed, $[]$ disappears.

Unlike BNF, the “language” of the metavariable/non-terminal A (the set of strings derived from A using roughly the rules of BNF plus hole filling) is not context-free and so MBNF certainly isn't.

²We use metavariable to mean a variable at the meta-level which denotes something at an object-level.

³The root node is on the spine. If A is applied to B by an application on the spine, the root node of A is on the spine and the root node of B is not. If a node on the spine is an abstraction each of its children is on the spine.

⁴A balanced segment is one where each application has a matching abstraction and where each application/abstraction pair contains a balanced segment.

MBNF Mixes Math Stuff With BNF-Style Notation Germane and Might [4, pg 20] mix BNF-style notation freely with mathematical notation in such a way that the resulting grammar relies upon both sets produced from the result of MBNF calculations and MBNF production rules which use metavariables defined using mathematical notation:

$$\begin{array}{ll}
 u \in UVar = \text{a set of identifiers} & ccall \in CCall ::= (q e^*)_\gamma \\
 k \in CVar = \text{a set of identifiers} & e, f \in UExp = UVar + ULam \\
 lam \in Lam = ULam + CLam & q \in CExp = CVar + CLam \\
 ulam \in ULam ::= (\lambda e(u^* k) call) & \ell \in ULab = \text{a set of labels} \\
 clam \in CLam ::= (\lambda_\gamma(u^*) call) & \gamma \in CLab = \text{a set of labels} \\
 call \in Call = UCall + CCall & \\
 ucall \in UCall ::= (f e^* q)_\ell &
 \end{array}$$

The results of math computations are interleaved with MBNF production rules, not just applied after the results of the production rules have been obtained. This grammar uses $\bullet_1 \in \bullet_2$ to mean “ $\bullet_2$ is the language of $\bullet_1$ ” (this is the case in both the MBNF production rules ($::=$) and the math itself ($=$)).

MBNF Has at Least the Power of Indexed Grammars Inoe and Taha [7, pg 361] use this MBNF:

$$\mathcal{E}^{\ell, m} \in ECtx_{\mathbf{n}}^{\ell, m} ::= \dots | \langle \mathcal{E}^{\ell+1, m} \rangle | \dots$$

This suggests that MBNF deals with the family of indexed grammars [6, p 389-390], which is yet another reason it’s not context-free. The $\ell + 1$ is a calculation that is not intended to be part of the syntax. The production rule above defines an infinite set of metavariables ranging over different sets.

MBNF Allows Arbitrary Side Conditions on Production Rules An example of a production rule with a side condition can be found in Chang and Felleisen [2, p 134]:

$$E = [] \mid Ee \mid A[E] \mid \hat{A}[A[\lambda x. \check{A}[E[x]]]E] \quad \text{where } \hat{A}[\check{A}] \in A$$

It is possible to make side conditions that prevent MBNF rules from having a solution. A definition for MBNF can help in finding restrictions on side conditions that ensure MBNF rules actually define something.

MBNF “Syntax” Can Contain Very Large Infinite Sets Toronto and McCarthy [23, p 297] use the following MBNF:

$$e ::= \dots | \langle t_{set}, \{e^{*\kappa}\} \rangle$$

We are told $\{e^{*\kappa}\}$ denotes “sets comprised of no more than κ terms from the language of e ”. The author does not state what κ is, but elsewhere in the paper it is an inaccessible cardinal. It seems as though κ is also intended to be an inaccessible cardinal here. This section of an MBNF for e is taken from a larger MBNF that contains a term which ranges over all the encodings of all the hereditarily accessible sets. BNF, by contrast, only deals with strings of finite length.

MBNF Allows Infinitary Operators Fdo, Díaz and Núñez [12, p 539] write an MBNF with the following operator:

$$P ::= \dots \mid \prod_{i \in I} P_i \mid \dots$$

The authors state this is infinitary (i.e. we should regard I to be infinite). The authors tell us the MBNF this is taken from is defined by regarding (M)BNF expressions as fixed point equations and a least fixed point can be found by bounding the size of the possible set of indices by some infinite cardinal.

We may think of infinitary operators as allowing us to define trees of infinite breadth (i.e. trees whose internal nodes may have infinitely many direct children), where BNF only deals with finite strings.

MBNF Allows Co-Inductive Definitions Eberhart, Hirschowitz and Seiller [3, p 94] intend the following MBNF to define infinite terms co-inductively:

$$\begin{aligned} P, Q &::= \sum_{i \in n} G_i \mid (P \mid Q) \\ G &::= \overline{a}\langle b \rangle.P \mid a(b).P \mid \nu a.P \mid \tau.P \mid \heartsuit.P \end{aligned}$$

We may think of co-inductive definitions as allowing us to define trees of infinite depth (i.e. trees in which paths may pass through infinitely many nodes), where BNF only deals with finite strings.

2 A Method to Allow Reading Some Uses of Mathematical “Syntax”

This section defines *syntactic math text* (SMT) which will allow reading some uses of math text as being “syntax” and standing for essentially themselves, e.g., $1 + 3$ can continue to stand for 4 while $\lambda x.x$ can in some sense stand for itself. SMT plus a definition of the $::=$ notation allows us to interpret the more common uses of MBNF as they are written. It also provides some support for more complicated uses with a little extra machinery. We do not aim to cover every use of MBNF in the literature, but we hope to provide a good foundation which can be built upon.

As well as dealing with some of MBNF, SMT provides a more general notion of *objects* appearing within syntax that behave like equivalences over chunks of math-text representing syntax. This enables us to interpret working modulo equivalences on math-text representing syntax.

Kamareddine et al. [10] make the point that converting mathematical text to a form where it can be checked by a proof assistant is a process that involves both human input and intermediary translations. Our proposal focuses on the translation, performed by the reader, of math-text used to define syntax, as it appears in a document, to a more formal structure, which is not encoded in the language of a theorem prover or proof assistant.

Our proposal relies as much as possible on the mathematical meta-level. For example, we use ellipses and related methods for abbreviating sequences from the mathematical meta-level. Incomplete definitions (relying on some choice of metavariable) cause the resulting grammar to be defined as the output of a function depending on this choice. Any otherwise pointless statement of the form $x \in S$ declares x and any *decorated* x (e.g., $x_1, x_2, \dots, x', x''$, etc.) as a variable ranging over S .

Our proposal is intended to be descriptive rather than prescriptive. We aim to handle both historical documents and new works. For published uses of MBNF that our proposal fails to handle, this is a problem to be solved in future work. We do not aim at displacing the input languages of proof assistants or syntactic variants of BNF which already have solid definitions.

2.1 Objects, Arrangements, and Symbols

We now define the main notion of syntactic *objects* and the auxiliary notion of *arrangements*. In essence, syntactic objects are arrangements of *symbols*, numbers, and pointers to subobjects, where the arrangement can include left-to-right sequencing, superscripting, subscripting and overlining. We use pointers to subobjects inside objects rather than the subobjects themselves, because the sets within the model for objects would be too large otherwise and because we wanted to allow for objects to be nested within themselves, provided some syntax is added as part of this nesting. To support α -conversion and operators that are associative, commutative, idempotent, etc., the objects are defined so that in effect they work modulo an equivalence relation on arrangements that is defined separately.

Let s range over the set **Symbol** containing syntactic *symbols* to be used in arrangements. We require that **Symbol** is disjoint from all other sets defined here. We also require that some symbols are *not* in **Symbol**, namely the square brackets (“[” and “]”) and the special square symbol $\square$ (which represents a hole in which an object can be placed). The symbols can include letters, parentheses and other parenthesis-like symbols (e.g., $\langle$ and $\rangle$ and $\llbracket$ and $\rrbracket$), punctuation, and other symbols. Letters (Roman or Greek) used as syntactic symbols will be typeset using an upright sans-serif font to distinguish them from metavariables which are written in a slanted serif font (generally italics). For example, $\mathbf{a}$, $\mathbf{C}$, λ , and Γ could be syntactic symbols while a , C , λ , and Γ would be metavariables. We avoid using any particular letter both ways, except for symbols used in *names*, where for example x_i could be a syntactic name at the same time as x could be a metavariable ranging over names (see section 2.5).

The set **Object** of syntactic *objects* and the set **Arrangement** of syntactic *arrangements* are defined simultaneously. Let O range over **Object** and let A range over **Arrangement**. We represent each object by a member of the set **Pointer**. Let P_O be the pointer that indicates O . Let $\approx \subset \mathbf{Arrangement} \times \mathbf{Arrangement}$ be an equivalence relation that is reflexive on **Arrangement**. We require that if $A_1 \approx A_2$ and $A_1 \neq A_2$, then neither A_1 nor A_2 may have used the special object $\square$ in their construction.

The sets **Object** and **Arrangement** are the smallest sets satisfying the following conditions.

1. The *empty arrangement* ϵ is in **Arrangement**.
2. The core items of arrangements are symbols, pointers to objects, numbers, and overlined arrangements. For any symbol s , pointer P_O , number $n \in \mathbb{N}$, and nonempty arrangement $A \neq \epsilon$, all of the following are in **Arrangement**: s , P_O , n , and $\overline{A}$. Furthermore, these are all *core arrangements*, which are ranged over by the metavariable $\hat{A}$.
3. Left-to-right sequencing allows appending additional core arrangements to a non-empty arrangement. For any arrangement $A \neq \epsilon$ and core arrangement $\hat{A}$, it holds that $A\hat{A}$ is in **Arrangement**.
4. Superscripting, subscripting etc. are supported. For non-empty arrangements A , A_1 and A_2 , all of the following are in **Arrangement**: A^{A_1} , A_{A_2} , $A_{A_2}^{A_1}$, $A_1 A_2 \dots$
5. If S contains does not contain any arrangements consisting of a bare pointer to an object, S is non-empty, and $|S| \leq \aleph_0$, then $S \in \mathbf{Object}$.

If $S \subset \mathbf{Arrangement}$ is not an equivalence class of $\approx$ or any members of S are ill-formed, then S is *ill-formed*. An arrangement is ill-formed iff any of its subcomponents is ill-formed. (Symbols and natural numbers are well formed.)

(Thus, it is allowed to build an object from ill-formed arrangements, and the resulting object is ill-formed.)

6. There is a special symbol in **Object** indicating a *hole* $\square$ in which an object is to be placed.

There are various reasons why we have built equivalence classes into arrangements rather than making them identical to math-text. We want to eventually support math stuff in syntax, with math stuff containing objects not arrangements. We want to allow object-to-object operations in production rules. When we define equivalences inductively over arrangements we want some of that structure to be represented by our model.

We write $[A]_{\approx}$ for the object that contains all the arrangements equivalent to A by the equivalence relation $\approx$. Only objects of the form $[A]_{\approx}$ are *well formed*.

2.2 Syntax Shorthand: Arrangement Coercions

From example 2.13, the reader will observe that it's cumbersome to write P_O in so many places when all we're interested in is the identity for objects. We introduce the following convention:

Convention 2.1 (Coercing Objects to Pointers). We allow O to be written instead of P_O in an arrangement. □

EXAMPLE 2.2. The expression $[\lambda O_1. O_2]_{\approx}$, stands for $[\lambda P_{O_1}. P_{O_2}]_{\approx}$. □

We define *meta-level parentheses* to be those parentheses which surround a single object and which may optionally be omitted from some arrangements with a similar form.⁵

It is still cumbersome to write $[\cdot]_{\approx}$ in so many places. One of the ways we deal with this is to arrange for this to happen automatically at places where a piece of meta-level syntax requires an arrangement to be regarded as an object.

Convention 2.3 (Coercing Arrangements to Objects). We require that when an arrangement A is written, but the surrounding context only makes sense if the value of the expression is an object, then the arrangement A is implicitly coerced to the object $[A]_{\approx}$, as though the latter had been written instead. As a special case of this, we require that an arrangement that containing meta-level parentheses is to be read as though the parentheses were instead a use of $[\cdot]_{\approx}$. □

Convention 2.4 (Coercing Arrangements to Pointers). We require that when an arrangement A is written, but the surrounding context only makes sense if the value of the expression is a pointer, then the arrangement A is implicitly coerced to the pointer to the object given by convention 2.3. □

Due to the combination of convention 2.3, convention 2.4 and the tight restrictions on where round parentheses can occur in proper arrangements, most uses of round parentheses will not be symbols that are part of syntactic arrangements but instead will be part of the meta-level mathematical reasoning.

EXAMPLE 2.5. The expression $(O_1 O_2) O_3$, which contains meta-level parentheses, stands for $[O_1 O_2]_{\approx} O_3$. If we write $O = (O_1 O_2) O_3$, then this stands for writing $O = [[O_1 O_2]_{\approx} O_3]_{\approx}$, because the equation's left-hand side must be an object due to the declaration that the metavariable O ranges over Object. □

We have left $\approx$ mostly unspecified so far. The sets Object and Arrangement do not depend on $\approx$, but their subsets of well formed objects and arrangements do depend on $\approx$. The definition of $\approx$ may be adjusted by the authors of a paper at any point, and the set of well formed objects in scope will therefore change at the times these adjustments are made. The effect of convention 2.3 will similarly change; the same expression can denote different objects at different places if there is an intervening change to $\approx$.

⁵We largely leave it up to the reader to determine which parentheses are meta-level. If a primitive constructor (section 2.4) appears inside some arrangements with parentheses surrounding it and other arrangements without them, it usually indicates these parentheses are meta-level. Similarly, parentheses which only surround a single metavariable corresponding to an object are frequently meta-level. Parentheses surrounding syntax which is to be thought of as a sequence are normally not meta-level. To help with this ambiguity, from this point forward all parentheses appearing in arrangements inside this document are meta-level.

2.3 Contexts and Hole Filling

A *context* is an object $O \in \text{Object}$ with at least one use of the special hole object $\square$. The number of hole symbols in an object or arrangement is its *arity*. We now define *context-hole filling* for arbitrary objects and arrangements (although it will in general only do something useful for well formed objects and arrangements with the correct arity). Let the operations $O[O_1, \dots, O_n]$, $P_O[O_1, \dots, O_n]$ and $A[O_1, \dots, O_n]$ which fill the holes reachable from O , P_O and A with the objects in the sequence $\vec{O} = [O_1, \dots, O_n]$ be defined as follows:

1. $P_O \vec{O} = P_{O'}$ and $O \vec{O} = O'$ iff $\text{fill}(O, \vec{O}) = (O', [])$. Similarly, $A \vec{O} = A'$ iff $\text{fill}(A, \vec{O}) = (A', [])$. The results of $\text{fill}(O, \vec{O})$ and $\text{fill}(A, \vec{O})$ are undefined except where explicitly defined below. (The result is undefined unless all of the replacements are used, so the number of replacements must match the arity.)
2. $\text{fill}(\square, [O] \cdot \vec{O}) = (O, \vec{O})$. (Each hole uses up one of the replacements.)
3. $\text{fill}(\{A\}, \vec{O}) = ([A']_{\approx}, \vec{O}')$ if $\text{fill}(A, \vec{O}) = (A', \vec{O}')$. (Context-hole filling in a well formed context can only descend inside an arrangement that is alone in its equivalence class. This is part of the motivation for our requirement that $\approx$ must not relate distinct arrangements containing holes.)
4. $\text{fill}(O, \vec{O}) = (O, \vec{O})$ if O is not a context. (This is the only way context-hole filling can skip over embedded objects which are non-singleton equivalence classes of arrangements.)
5. $\text{fill}(s, \vec{O}) = (s, \vec{O})$ and $\text{fill}(n, \vec{O}) = (n, \vec{O})$.
6. Context-hole filling essentially traverses the arrangement tree in a left-to-right order filling in holes in the order it encounters them. Thus, for any arrangements A , A_1 , and A_2 , core arrangement $\hat{A}$, and object sequences $\vec{O}_1$, $\vec{O}_2$, $\vec{O}_3$, and $\vec{O}_4$, if it holds that

$$\begin{array}{ll} \text{fill}(A, \vec{O}_1) = (A', \vec{O}_2) & \text{fill}(A_1, \vec{O}_2) = (A'_1, \vec{O}_3) \\ \text{fill}(A_2, \vec{O}_3) = (A'_2, \vec{O}_4) & \text{fill}(\hat{A}, \vec{O}_2) = (\hat{A}', \vec{O}_3) \end{array}$$

then all of these must follow:

$$\begin{array}{ll} \text{fill}(A\hat{A}, \vec{O}_1) = (A'\hat{A}', \vec{O}_3) & \text{fill}(A_{A_2}^{A_1}, \vec{O}_1) = (A_{A_2}^{A'_1}, \vec{O}_4) \\ \text{fill}(A^{A_1}, \vec{O}_1) = (A^{A'_1}, \vec{O}_3) & \text{fill}(A_{A_1}, \vec{O}_1) = (A_{A_1}', \vec{O}_3) \\ \text{fill}(\overline{A}, \vec{O}_1) = (\overline{A'}, \vec{O}_2) & \text{fill}(\underline{A}, \vec{O}_1) = (\underline{A'}, \vec{O}_2) \end{array}$$

7. $\text{fill}(P_O, \vec{O}_1 \cdot \vec{O}_2) = \text{fill}(P_{O_3}, \vec{O}_2)$ if $\text{fill}(O, \vec{O}_1 \cdot \vec{O}_2) = \text{fill}(O_3, \vec{O}_2)$. (Context-hole filling descends object pointers until it encounters a hole)

EXAMPLE 2.6. Here are some examples of context-hole filling :

$$\begin{array}{ll} (\square \square)[O, O] = OO & \square[O] = \underline{O} \\ (\square \rightarrow O_1)[O_2 \rightarrow O_2] = (O_2 \rightarrow O_2) \rightarrow O_1 & (\square := \square, \square)[O_1, O_2, O_3] = (O_1 := O_2, O_3) \end{array} \quad \square$$

We now will define (S_1, S_2) -Context to be the the contexts which act as functions from S_1 to S_2 , i.e., the set of every context O_c of arity 1 such that for all $O \in S_1$ it holds that $O_c[O] \in S_2$. Let $S\text{-Context} = (S, S)\text{-Context}$.

Given a relation $\mathcal{R}$ such that $(\text{domain}(\mathcal{R}) \cup \text{range}(\mathcal{R})) \subseteq S \subseteq \text{Object}$, let $[\mathcal{R}]^S$ denote the S -compatible closure of $\mathcal{R}$, defined as follows: if $O_c \in S\text{-Context}$ and $O_1 \xrightarrow{\mathcal{R}} O_2$,⁶ then $O_c[O_1] \xrightarrow{[\mathcal{R}]^S} O_c[O_2]$. Let $[\mathcal{R}]$ denote $[\mathcal{R}]^S$ for some set S which the reader can infer from the context of discussion.

Let c range over *primitive constructors*, non-hole objects whose only immediate subobjects are $\square$.

EXAMPLE 2.7. Here are some examples of primitive constructors [2, p 134], [22, p386], [7, pg 360], [16]:

$$(\square \square) \quad \square \downarrow \square \cdot \square \quad !\square \quad \langle \square \rangle \quad \square + \square \quad \square = \square \in \square \quad \square$$

Every well formed non-hole object O can be decomposed into a primitive constructor and the subobjects to be placed in the primitive constructor's holes. A *primitive constructor decomposition* of O is a pair $(c, \vec{O})$ such that $O = c\vec{O}$. An object will have one primitive constructor decomposition for each of its arrangements. Furthermore, the subobjects in a decomposition can be recursively decomposed similarly. A recursive decomposition of an object into primitive constructors is very similar to the concept of an *abstract syntax tree* of a string in a language defined by a grammar. If any of the equivalence classes in an object are non-singletons, then the object will not have a unique recursive decomposition.

EXAMPLE 2.8. Some examples of recursive decomposition of an object into primitive constructors can already be seen in example 2.6. Here are some additional examples:

$$\langle (!O) \rangle = \langle \square \rangle [!\square[O]] \quad (O_1 + O_2) + O_3 = (\square + \square) [(\square + \square)[O_1, O_2], O_3]$$

□

2.4 Syntax Shorthand: Primitive Constructor Decomposition

Convention 2.3 allows avoiding the need to write $[\cdot]_{\approx}$ by implicitly invoking $[\cdot]_{\approx}$ at obvious arrangement boundaries and also at most uses of $(\cdot)$ in arrangements. For example, convention 2.3 allows us to know that the expression $(O_1 @ O_2) @ O_3$ stands for the object whose primitive constructor decomposition is given by $O' = c_{@}[c_{@}[O_1, O_2], O_3]$ where $c_{@} = \square @ \square$.

But the shorthand notation provided by convention 2.3 is not enough. Additionally, we want to allow inferring uses of $[\cdot]_{\approx}$ in other places in the middle of what appear to be arrangements. As a concrete example, we want to allow inferring that the expression $O_1 @ O_2 @ O_3$ stands for the same object as the expression $(O_1 @ O_2) @ O_3$, namely the object O' mentioned in the previous paragraph. We want that the expression $O_1 @ O_2 @ O_3$ must *not* stand for the object whose primitive constructor decomposition is $c''[O_1, O_2, O_3]$ where $c'' = (\square @ \square @ \square)$.

To provide the additional shorthand notation that is needed, we establish mechanisms for (1) declaring primitive constructors and (2) parsing arrangements. We build the parsing mechanism by adapting the notions of operator precedence and declared associativity from parsing of languages to our setting; this will allow splitting what appears to be a single primitive constructor into multiple primitive constructors.

As an auxiliary device, we define splicing of arrangements. Remember that every arrangement is, in effect, a sequence of core arrangements (symbols, objects, numbers, or overlined arrangements), possibly superscripted or subscripted. An arrangement A' can be *spliced* into another arrangement A'' by inserting the main core arrangement sequence of A' into one of the core arrangement sequences of A'' in place of an occurrence of $\square$.

⁶ $O_1 \xrightarrow{\mathcal{R}} O_2$ is an alternative notation for $(O_1, O_2) \in \mathcal{R}$. See appendix A.4 for details.

Convention 2.9 (Declaring and Parsing Primitive Constructors).

1. Unless prevented by part 2 of this convention, at the first use of a proper arrangement A , if there is a primitive constructor $c = \{A'\}$ and objects $O_1, \dots, O_n$ such that $\{A\} = c[O_1, \dots, O_n]$, then this use of A *declares* the primitive constructor c and the arrangement A' . Note that A' differs from A exactly in having $\square$ in place of every non- $\square$ object appearing in A .
2. At each place where we coerce an arrangement A into an object O using convention 2.3, the arrangement A is inspected to see if it can be built by splicing together already-declared arrangements. If A can be built entirely by splicing together already-declared arrangements, and then filling the holes in the splicing result with objects, and there is no explicit indication forbidding the use of this convention, then A is to be interpreted as though it had been written with uses of $[\cdot]_{\approx}$ around each splice point. If there is more than one way A can be built by splicing already-declared arrangements, then it must be specified somewhere which one to choose. (This choice will typically involve notions of operator precedence and declarations of associativity.) $\square$

EXAMPLE 2.10. Suppose we have written the expressions $\langle O_1 \rangle$ and $!O_2$. This declares the primitive constructors $\langle O_1 \rangle$ and $!O_2$. If we then write $O = \langle !O' \rangle$, then by convention 2.9 this produces the same result as writing $O = \langle (!O') \rangle$. This happens because the arrangement $\langle !O' \rangle$ can be built by splicing $! \square$ into $\langle O_1 \rangle$ and then filling the hole with O' .

(If we wanted to avoid the interpretation of convention 2.9, we could do so by avoiding the implicit coercion of convention 2.3 and writing instead $O = [\langle !O' \rangle]_{\approx}$, which would use the primitive constructor $\langle !\square \rangle$ instead of the two smaller primitive constructors $\langle \square \rangle$ and $!\square$.)

Suppose we write the expression $O_1 @ O_2$. This declares the primitive constructor $c_{@} = \square @ \square$. If we then state that $c_{@}$ is left-associative, then writing $O = O_1 @ O_2 @ O_3$ produces the same result as writing $O = (O_1 @ O_2) @ O_3$. If we did not give the associativity of $c_{@}$, then writing $O = O_1 @ O_2 @ O_3$ would be an error, because there are multiple distinct ways the arrangement $\square @ \square @ \square$ can be built by splicing the arrangement $\square @ \square$ into itself. $\square$

2.5 Names, Binding, α -Conversion, and Substitution

The relation $\approx$ provides a mechanism for working with syntax considered modulo equivalences on arrangements. One of the most important equivalences is the notion of α -conversion which renames *bound names*.⁷

Some of the members of **Object** can be declared to be *names*. The names may be furthermore subdivided into groups. Formally, the concepts of names and groups of names are given by an equivalence relation $\sim \subset \mathbf{Object} \times \mathbf{Object}$ which relates names in the same group. An object O is a *name* iff $O \sim O$. Declaring a subset $S \subset O$ to be a *name group* is the same as declaring that S is a $\sim$ -equivalence class. The definition of $\sim$ will be extended incrementally with declarations of groups. Any objects that have not been declared to be related by $\sim$ are *not* related by $\sim$. To keep things simple we require that no name contains another name (of the same group or of a different group) as a subobject.

Specific primitive constructors can be declared to *bind* a name placed in one of the constructor's holes across some of the constructor's holes. We define the *free names* of an object O , written $\text{FN}(O)$:

1. If O is a name, then $\text{FN}(O) = \{O\}$.

⁷ We do not give an especially sophisticated notion of binding here. We are only interested in providing a concept of binding that can be readily grasped and is sufficiently general for wide use in a variety of grammars. The notion of equivalence we provide is intended to be used in defining other syntactic equivalences in addition to α -equivalence.

2. Otherwise, if $\text{FN}(O)$ is defined, it is as follows.

First, we must define the free names of primitive constructor decompositions (p.c.d.'s) of O . Suppose $O = c[O_1, \dots, O_n]$ gives one such p.c.d. Let $\mathcal{S}_i$ be the names bound by c in O_i for $1 \leq i \leq n$. Then $\text{FN}(c, [O_1, \dots, O_n]) = \bigcup_{i \in \{1, \dots, n\}} \text{FN}(O_i) \setminus \mathcal{S}_i$.

If there exists a set $\mathcal{S}$ such that $\mathcal{S} = \text{FN}(c, \vec{O})$ for every p.c.d. $(c, \vec{O})$ of O , then $\text{FN}(O) = \mathcal{S}$.

The free names of an arrangement A are defined by $\text{FN}(A) = \text{FN}(\{A\})$. A name that is not free is *bound*.

EXAMPLE 2.11. Consider $c_\lambda = \lambda \square. \square$ of arity 2. Suppose we declare that c_λ binds any name placed in its first hole in both of its holes. Suppose we declare that $\{x_i \mid i \in \mathbb{N}\}$ is a name group. (We will in fact make both of these declarations later, so this example is not just hypothetical.) Suppose that we have not declared any bindings for the constructor $c_{@} = \square @ \square$. Then $\text{FN}((\lambda x_1. (x_1 @ x_2)) @ x_3) = \{x_2, x_3\}$.

Consider $c_{\text{let}} = (\text{let } \square = \square \text{ in } \square)$ of arity 3. Suppose we declare that c_{let} binds any name placed in its 1st hole in its 1st and 3rd hole. Then $\text{FN}(\text{let } x_1 = x_3 \text{ in } (x_1 @ x_2)) = \{x_2, x_3\}$. $\square$

We now define the auxiliary notion of *name swapping*. Given two names O_x and O_y such that $O_x \sim O_y$, let $\text{swap}(O_x, O_y, O)$ be the object O' that results from replacing every occurrence of O_x in O by O_y , and *vice versa*. Let $\text{swap}(O_x, O_y, A)$ be defined similarly.

We now define α -conversion. Let $\equiv_\alpha$ be the smallest equivalence relation satisfying the following condition. For all O_x, O_y, O , and A , if $O_x \sim O_y$ and $\{O_x, O_y\} \cap \text{FN}(O) = \{O_x, O_y\} \cap \text{FN}(A) = \emptyset$, then $O \equiv_\alpha \text{swap}(O_x, O_y, O)$ and $A \equiv_\alpha \text{swap}(O_x, O_y, A)$.

Definition 2.12 (α -Conversion as a Syntactic Equivalence). If a paper says that it is “working modulo α ” or “identifying α -equivalent terms” that means $\equiv_\alpha$ restricted to arrangements is a subset of $\approx$, i.e., if $A_1 \equiv_\alpha A_2$ then $A_1 \approx A_2$. $\square$

Definition 2.12 implies that $\approx$ will change whenever adjustments are made to the declared bindings of primitive constructors or to the definition of $\sim$.

We now define the *substitution* operation, written as $O[O_x := O']$. This expression will be defined to stand for the result of replacing all free occurrences of O_x in O by O' . This operation must be defined carefully. The result of $O[O_x := O']$ must not allow names that are free in O' to be captured by bindings in O . Also, the operation must respect $\approx$ so that if both O and O' are well formed, then $O[O_x := O']$ is also well formed. Given a name O_x , define $O[O_x := O']$ formally as follows.

1. If $O = O_x$, then $O[O_x := O'] = O'$.
2. Otherwise, $O[O_x := O']$ is defined as follows.

First, we must define substitution for primitive constructor decompositions (p.c.d.'s). Given $O = c[O_1, \dots, O_n]$, let $\mathcal{S}$ be the subset of $\{O_1, \dots, O_n\}$ of names bound by this occurrence of c . If $\mathcal{S} \cap \text{FN}(O') \neq \emptyset$, then let $(c, [O_1, \dots, O_n])[O_x := O']$ be undefined.⁸ Otherwise, let $(c, [O_1, \dots, O_n])[O_x := O'] = c[O_1[O_x := O'], \dots, O_n[O_x := O']]$.

If there exists an O'' such that $O'' = (c, \vec{O})[O_x := O']$ for every p.c.d. $(c, \vec{O})$ of O such that $(c, \vec{O})[O_x := O']$ is defined, then $O[O_x := O'] = O''$. Otherwise $O[O_x := O']$ is undefined.⁹

EXAMPLE 2.13. Below, on the left are some example syntactic objects [2, p 134], [16], [22, p 386]. These objects may not be well formed, because the singleton sets may not be equivalence classes of $\approx$. The

⁸For simplicity, we do not check whether the substitution needs only to proceed into holes of c which are not subject to its bindings. This will behave well enough for our uses provided each group of names is big enough that fresh names can be found.

⁹So the substitution must be defined for at least one of the primitive constructor decompositions to get a defined result.

objects to the right of them are adjusted to be well formed (assuming the subobjects O_1 , to O_4 are well formed):

$$\begin{array}{ll} \{\lambda P_{O_1}.P_{O_2}\} & [\lambda P_{O_1}.P_{O_2}]_{\approx} \\ \{\Pi P_{O_1} : P_{O_2}.P_{O_3}\} & [\Pi P_{O_1} : P_{O_2}.P_{O_3}]_{\approx} \\ \{^{P_{O_1}} \downarrow \{P_{O_2} P_{O_3}\} \cdot P_{O_4}\} & [^{P_{O_1}} \downarrow [P_{O_2} P_{O_3}]_{\approx} \cdot P_{O_4}]_{\approx} \end{array}$$

□

2.6 Production Rules for Defining Syntactic Sets

We have already defined syntactic objects, but the set **Object** is too big. Carefully defined subsets of **Object** may be defined via *syntax production rules*, which we write in the form

$$\nu_1, \dots, \nu_n \in \mathcal{S} ::= \mathcal{A}_1 \mid \dots \mid \mathcal{A}_m$$

where $\nu_1, \dots, \nu_n$ are metavariables, $\mathcal{S}$ is the name of the subset of **Object** being defined, and $\mathcal{A}_1, \dots, \mathcal{A}_m$ are *alternatives*. Each alternative is either the special notation “ $\dots$ ” or an expression e , together with an optional side condition c (written “ e if c ”, where c is a formula containing only expressions), which evaluates to a member of **Object** when values are supplied for metavariables occurring in e , provided both c holds of that choice of metavariables. One can omit the “ $\in \mathcal{S}$ ”, allowing the reader to fill in $\mathcal{S}$ whose name is distinct from the names of all other declared sets. One can omit the side condition in which case we can read it as if true. One can provide a global side condition if c' which we read as appending $\wedge c'$ to all $\mathcal{A}$.

Such a syntax production rule has the following effects:

1. It declares $\mathcal{S}$ to be a set of syntactic objects, in particular the smallest one that satisfies all other constraints placed on it not just by this rule but also by the rest of the document.
2. It declares the metavariables $\nu_1, \dots, \nu_n$ to range over the set $\mathcal{S}$.
3. A global side condition if c' appends $\wedge c'$ to each $\mathcal{A}_1, \dots, \mathcal{A}_n$.
4. If each $\mathcal{A}_1, \dots, \mathcal{A}_n$ contains only undecorated instances of ν , then for any $\mathcal{A}$ containing multiple instances of ν and no side conditions containing ν that apply to $\mathcal{A}$, we can rewrite it with each ν given a different decoration. I.e., $m \in M ::= x \mid m m$ becomes $m, m_1, m_2 \in M ::= x \mid m_1 m_2$.
5. For each alternative $\mathcal{A}$ in the rule which is not “ $\dots$ ”, a constraint on the membership of $\mathcal{S}$ is added. The constraint is that for each legal choice¹⁰ of values for the metavariables occurring in $\mathcal{A}$, if O is the result of evaluating the expression e in $\mathcal{A}$ using those metavariable assignments, then $O \in \mathcal{S}$. Metavariables occurring in an alternative $\mathcal{A}$ that are not yet declared to range over any set are presumed to range over a countable set of objects disjoint from all the other sets of objects in the paper. This assumption is dropped if a value for a metavariable gets declared later in the paper and values for $\mathcal{A}$ are recalculated accordingly.
6. If the first alternative is not the special alternative “ $\dots$ ”, then any constraints on the membership of $\mathcal{S}$ established by earlier rules are forgotten.

¹⁰ By legal choice we mean a choice of metavariables matching the sets they are declared to range over and fulfilling any constraints added by any side conditions.

7. The rule triggers a recalculation of *all* of the sets declared by all syntax production rules. Such a recalculation is also triggered whenever the definition of $\approx$ is altered. Or when a definition of what metavariables range over is altered.

This recalculation evaluates all of the constraint expressions for all syntactic sets using the *current* bindings for all metavariables, set names, the equivalence relation $\approx$, etc., and rebinds the set names to the recalculated values in the subsequent text.¹¹

Multiple rules can be given for the same set S . If a later rule for S begins with the special alternative “ $\dots$ ”, then its alternatives are combined with the alternatives already in force for S . Usually the alternatives of the later rule replace the previous alternatives if this is not the case. However, if the author uses a single alternative in each of their production rules, then they normally expect these to be combined as though they had used $\dots$. The special alternative “ $\dots$ ” used as the final alternative of a rule has no mathematical consequence and is used only as a signal to the reader warning that there will be later rules for the same set.

When a syntax alternative is intended to allow building terms from multiple subterms of the same set, it is necessary to use distinct metavariables for each possible subterm to allow the subterms to differ. It is always possible to find distinct metavariables for the same set by using subscripts.

EXAMPLE 2.14. We can define the usual *simple types* like this:

$$\begin{aligned} a, b \in \text{Ty-Variable} &::= a_i \\ T \in \text{Simple-Type} &::= a \mid T_1 \rightarrow T_2 \end{aligned}$$

Given this definition, a possible example type is $T_0 = a \rightarrow (b \rightarrow a)$. In this example T_0 , we leave unspecified which exact type variables are used. We could make T_0 concrete by specifying $a = a_0$ and $b = a_1$ yielding $T_0 = a_0 \rightarrow (a_1 \rightarrow a_0)$. If we had written the second alternative in the production rule for Simple-Type as $T \rightarrow T$, then the type T_0 would not be allowed and we could only write types like $a \rightarrow a$ and $(a \rightarrow a) \rightarrow (a \rightarrow a)$ where both arguments of each $\rightarrow$ are equal. $\square$

EXAMPLE 2.15. We can define the *lambda calculus* like this:

$$e \in \text{exp} ::= v \mid \lambda v. e \mid e e$$

Each v ranges over a countable set of objects disjoint from the objects produced by the other production rules. The production rule $e \in \text{exp} ::= v$ can be read as giving us the constraint $\text{var} \subseteq \text{exp}$. The constraint $\{[\lambda P_v. P_e]_{\approx} \mid \text{ptr}(v) = P_v \wedge \text{ptr}(e) = P_e\} \subseteq \text{exp}$ is given by $e \in \text{exp} ::= \lambda x. e$. The constraint $\{[P_{e_1} P_{e_2}]_{\approx} \mid \text{ptr}(e_1) = P_{e_1} \wedge \text{ptr}(e_2) = P_{e_2}\} \subseteq \text{exp}$ is given by $e \in \text{exp} ::= e e$. We pick the least $\text{exp} \subseteq \text{Object}$ and $\text{var} \subseteq \text{Object}$ satisfying these constraints with an ordering given by the subset relation.

In addition to declaring e as ranging over exp this definition also declares $e_1, e_2, \dots, e', e''$ etc. to range over exp and similarly for $v \in \text{var}$. The subset of Object picked out by these constraints depends on the choice of equivalence relation $\approx$, in the lambda calculus this is most likely α equivalence, although it may also be the identity relation on Arrangement .

In order to be confident that this set can be picked out (e.g. for exp) we begin with $\text{exp}^0 = \emptyset$ and let exp^1 contain all the things exp must contain if exp is at least exp^0 and so on for each $+1$ case. For a limit point ε we let exp^ε be $\bigcup_{i=0}^{\varepsilon} \text{exp}^i$. We take the least fixed point of the function $f \in \mathcal{P}(\text{Object}) \rightarrow \mathcal{P}(\text{Object})$

¹¹It is an error if there is not a unique assignment of smallest values to the declared sets. Normally, the existence of a unique assignment will be provable using a fixed point theorem like the Knaster-Tarski theorem. However, the notation allows putting strange side conditions in the constraint expressions in alternatives, and this can cause a failure.

such that $f(\exp^i) = \exp^{i+1}$ over some appropriately large set of $\exp^i$ ordered but the subset relation (this is smaller than $\mathcal{P}(\text{Object})$).

We define the reduction relation as the e -compatible closure of the smallest β (in the ordering given by $\subseteq$) satisfying the constraint $(\lambda v.e_1)e_2 \xrightarrow{\beta} (e_1[v := e_2])$. The notation $O_1[O_2 := O_3]$ is defined in section 2.5. We do the same for $\lambda v.e_1 v \xrightarrow{\beta} e$. Note that because we have bracketed the term after substitution we are able to reapply equivalences that may have otherwise been lost in the process. $\square$

EXAMPLE 2.16. Given the definition of simple types in Example 2.14 we can define the *simply typed lambda calculus* as follows:

$$\hat{e} \in \text{texp} ::= v \mid \hat{e} \hat{e} \mid \lambda v : T. \hat{e}$$

$\square$

EXAMPLE 2.17. We can extend example 2.14 with records in a similar way to Pierce [15, pg 129]. We can define *lambda calculus with records* like this:

$$\begin{aligned} l \in \text{label} &::= y_i & R \in \text{Type-Records} &::= \epsilon \mid l : T, R & \text{ where } l \notin \text{lab}(R) \\ \hat{e} \in \text{texp} &::= \dots \mid \{r\} \mid \hat{e}.l \\ t \in \text{Record-Type} &::= T \mid \{R\} & r \in \text{Term-Records} &::= \epsilon \mid l = \hat{e}, r & \text{ where } l \notin \text{lab}(r) \end{aligned}$$

Where we define lab s.t. $\text{lab}(\epsilon) = \emptyset$, $\text{lab}(l : T, R) = \{l\} \cup \text{lab}(R)$ and $\text{lab}(l = \hat{e}, r) = \{l\} \cup \text{lab}(r)$. Both r and R are equivalent up to reordering (i.e. $l : T, R \approx R, l : T$ and $l = \hat{e}, r \approx r, l = \hat{e}$). Here, $\approx$ is the smallest equivalence relation fulfilling these constraints. It is defined incrementally over each R and each r as a new one is added.

We add a rewriting rule:

$$\{l = v, r\}.l \xrightarrow{\text{RCD}} v$$

For each $* \in \{x \in \text{Object} \times \text{Object} \mid x = \beta\} \cup \{x \in \text{Object} \times \text{Object} \mid x = \eta\} \cup \{\text{RCD}\}$ we add additional constraints:

$$\frac{(\hat{e}_1 \xrightarrow{*} \hat{e}_2)}{(\hat{e}_1.l \xrightarrow{*} \hat{e}_2.l)} \quad \frac{(\hat{e}_1 \xrightarrow{*} \hat{e}_2)}{(\{r, l = \hat{e}_1\} \xrightarrow{*} \{r, l = \hat{e}_2\})}$$

(The horizontal line is read as the logical operator $\Rightarrow$). $\square$

3 Model for Syntactic Math Text

In this section we show that there is a model for SMT. In order to do so, we choose sets to represent Symbol, Pos, Pointer, $\square$, ϵ and $\bar{B}$. Our invariant constraints are those that will hold of sets thought to approximate Object and Arrangement in the proof these are well defined. Our constraints on the final selection will only hold of the set we pick out from these approximations.

Definition 3.1 (Symbol, Pos, Pointer, $\{\square, \epsilon, \bar{B}\}$). We can create a countable set, D , representing symbols, accenting and positioning from the ordinals¹² following ω which are themselves smaller than 2ω . We pick a finite set of elements, Pos , from D to represent the positions subscript, superscript, pre-subscript, pre superscript, text above, text below etc (at least as many as positions as detailed in the OpenDocument [9] standard). We pick out an element of D which we call $\bar{B}$. We pick out a element D to represent the context-hole $\square$, and one to represent the empty arrangement ϵ . We let the remainder of the elements in D represent Symbol (at least as many symbols as in unicode). $\square$

¹²With the Von Neumann encoding

Definition 3.2 (Invariant Constraints).

$\text{ptr} \in \text{Object} \rightarrow \text{Pointer}$ and ptr is a bijection between Object and Pointer .

$\bar{B} \notin \text{Object} \wedge \bar{B} \notin \text{Arrangement}$.

$\square \in \text{Object} \wedge \square \notin \text{Arrangement} \wedge \epsilon \in \text{Arrangement}$

$\text{Pos} \perp \text{Object} \wedge \text{Pos} \perp \text{Arrangement}$.

$\text{Symbol} \subset \text{Core}$.

$\mathbb{N} \subset \text{Core}$.

$\text{Core} \subset \text{Arrangement}$. □

Definition 3.3 (Constraints on Final Selection).

$\text{Pointer} \subset \text{Core}$.

If $A \in \text{Arrangement}$, $A \neq \epsilon$ and $x \in \text{Symbol}$ then $(\bar{B}, x, A) \in \text{Core}$.

$\text{Arrangement} \times \text{Core} \subset \text{Arrangement}$.

If $A \in \text{Arrangement}$ and $x \in \text{Pos} \rightarrow \text{Arrangement} \setminus \{\epsilon\}$ and $x \neq \emptyset$ then

$(A, x) \in \text{Arrangement}$.

If $S \subset \text{Arrangement}$ and $|S| \leq \aleph_0$, then $S \in \text{Object}$.¹³ □

Theorem 3.4. *Object and Arrangement are well defined.*

Proof Sketch. We define a sequence of sets thought to contain closer approximations of Object and Arrangement until some member contains a model for Object and Arrangement themselves. The smallest set in our sequence contains all tuples of:

1. The set containing $\square$ (approximating Object).
2. An injective function $p \in \{\square\} \rightarrow \text{Pointer}$ (approximating ptr).
3. $\text{Symbol} \cup \mathbb{N} \cup \{\epsilon\}$ (approximating Arrangement).
4. $\text{Symbol} \cup \mathbb{N}$ (approximating Core).

Each subsequent set in our sequence contains those tuples of sets which would be added by applying our constraints as though Object were its approximation, Pointer were its approximation and Arrangement were its approximation. Where our sequence reaches a limit point each set in each tuple is calculated as though Arrangement was the union of arrangements up to that point (apart from the set approximating Arrangement which also gets the pointers to the approximation of Object at that limit).

These sets remain sufficiently small to pick mappings for Object . Further, there is a fixed point for the function mapping each member of this sequence to the member above it. From this fixed point we can select a model for Object and Arrangement .

(Full proof in Appendix B) □

4 How Can This Definition be Used?

Non-MBNF “Grammars” As well as covering some uses of MBNF to define syntax, SMT also provides us with a notion of what it means to use the structures of math-text together with syntactic equivalences, even in documents where MBNF does not feature, or where MBNF is mixed with other notation for picking out objects. Coverage of this sort may require users to select appropriate sets of objects that resolve ambiguities.

¹³We do not bother restricting objects to only include proper arrangements here as it does not particularly affect the logic of the proof. Provided one can pick out unique members from Symbol for left parenthesis, right parenthesis and comma, its not too hard to express what it means for an arrangement to be proper with a logical formula.

A Flexible Notion of Equivalence Not only is the notion of equivalence presented in SMT sufficient to deal with α -equivalence over finite terms, regardless of how binding may be represented in the syntax, it also deals with things like equality up to reordering of finitely many chunks of syntax and equality of finitely many compositions with zero, both of which appear in the π -calculus [1]. It deals with many of the equivalences an author might define using “=,” provided they do not quantify over an uncountable set when using it. Furthermore it provides tools to consider equalities over sub-objects, not just the syntactic objects themselves, which can be vital when talking about the structure of a grammar.

Combining Objects in Math-Text SMT deals with most combinations of characters likely to appear in math text used to represent “syntax” in a fairly general manner (it does not deal with matrices/grids, numbers other than the naturals, or an use of sets that cannot be thought of in terms of equivalences up to reordering and repetition on finite lists of elements, but none of these is likely to appear in “syntax”).

Automatic Bracketing Since SMT preserves the tree-like structure of syntax, it can be readily used for grammars where authors treat bracketing as optional. We also give authors the option of making this structure more explicit by primitive decomposition. Bracketing structures may often also be derived by noticing where objects appear in production rules.

Functionality Inherited From BNF Our definition extends the basic functions covered by BNF to MBNF and the richer syntactic structures that are represented by math-text. Substitution of non-terminals becomes assigning values to metavariables and choice of production rules remains supported.

Hole Filling The following chunk of the MBNF we took from Chang and Felleisen [2, p 134] defining A can be handled by our definition using convention 2.3:

$$\begin{aligned} e &= x_i \mid \lambda x_i. e \mid e e \\ A &= [] \mid A[\lambda x_i. A] e \end{aligned}$$

5 Related work

OTT [19] provides a formal language for writing specifications like those written in MBNF. The process of moving from an OTT specification to an MBNF can be performed automatically. However, the focus of this article is moving the other way — interpreting MBNF without requiring it to be specified in a theorem-prover friendly format. Furthermore, we wish to provide a general mathematical intuition suitable for translation to multiple theorem provers, whereas OTT focuses on translating to COQ 8.3, HOL 4 and Isabelle directly, but offers less support for those seeking a general mathematical intuition. In addition, OTT only supports context hole filling for contexts with a single hole and currently does not support rules being used coinductively. We already handle more cases of context hole filling and we aim to deal with coinduction, though SMT as it stands doesn’t.

Guy Steele [20] covers many of the notational variants of BNF, including some MBNFs. However, Steele’s focus is primarily on surface differences. He does not discuss how the underlying mathematical structure of MBNF differs wildly from BNF.

Grewe et al. [5] discuss the exploration of language specifications with first-order theorem provers. However, they still require the reader to be able to intuitively translate language specifications to a sufficiently formal language first. This is the part of language specification checking this paper aims to help with.

Reynolds [17, 1-51] has the best attempt at a definition of MBNF, which he calls “abstract syntax”,¹⁴ which we could find after looking through the books in our collection. However, he only deals context-free grammars and in many places he proceeds by example.

6 Future Work

While we do not deal with trees of infinite breadth or depth here, we hypothesise that the method outlined in this document could be used on trees with countably infinite breadth and depth. The main difference in doing so would be that Object and Arrangement would likely have to be of cardinality $\aleph_2$, rather than $\aleph_1$, but apart from that it seems likely a similar proof would work.

While we provide some powerful tools for writing syntax patterns more explicitly and dealing with numbers in the syntax, we do not provide procedures for generating countably many production rules. Guy Steele [20] has done work in this area, but doesn’t address differences between MBNF and BNF.

References

- [1] Marco Carbone, Kohei Honda & Nobuko Yoshida (2007): *Structured Communication-Centred Programming for Web Services*. In Rocco De Nicola, editor: *Programming Languages and Systems*, Springer Berlin Heidelberg, Berlin, Heidelberg, pp. 2–17.
- [2] Stephen Chang & Matthias Felleisen (2012): *The Call-by-need Lambda Calculus, Revisited*. In Seidl [18], pp. 128–147.
- [3] Clovis Eberhart, Tom Hirschowitz & Thomas Seiller (2015): *An Intensionally Fully-abstract Sheaf Model for π^** . In Lawrence S. Moss & Pawel Sobocinski, editors: *6th Conference on Algebra and Coalgebra in Computer Science (CALCO 2015)*, *Leibniz International Proceedings in Informatics (LIPIcs)* 35, Schloss Dagstuhl–Leibniz-Zentrum fuer Informatik, Dagstuhl, Germany, pp. 86–100, doi:10.4230/LIPIcs.CALCO.2015.86. Available at <http://drops.dagstuhl.de/opus/volltexte/2015/5528>.
- [4] Kimball Germane & Matthew Might (2017): *A Posteriori Environment Analysis with Pushdown Delta CFA*. In: *Proceedings of the 44th ACM SIGPLAN Symposium on Principles of Programming Languages*, ACM, New York, NY, USA.
- [5] Sylvia Grewe, Sebastian Erdweg, André Pacak, Michael Raulf & Mira Mezini (2018): *Exploration of language specifications by compilation to first-order logic*. *Sci. Comput. Program.* 155, pp. 146–172, doi:10.1016/j.scico.2017.08.001. Available at <https://doi.org/10.1016/j.scico.2017.08.001>.
- [6] John E. Hopcroft, Rajeev Motwani & Jeffrey D. Ullman (2006): *Introduction to Automata Theory, Languages, and Computation (3rd Edition)*. Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA.
- [7] Jun Inoue & Walid Taha (2012): *Reasoning About Multi-stage Programs*. In Seidl [18].
- [8] Patrick D F Ion, Nico Poppelier, David Carlisle & Robert R Miner (2001): *Mathematical Markup Language (MathML) Version 2.0*. W3C Recommendation, W3C. <https://www.w3.org/TR/MathML2/chapter3.html>.
- [9] (2015): *Information technology – Open Document Format for Office Applications (OpenDocument) v1.2 – Part 1: OpenDocument Schema*. Standard, International Organization for Standardization, Geneva, CH.
- [10] Fairouz Kamareddine, Joe Wells, Christoph Zengler & Henk Barendregt (2014): *Computerising Mathematical Text*. In Jörg H. Siekmann, editor: *Computational Logic, Handbook of the History of Logic* 9, North-Holland, pp. 343 – 396, doi:<https://doi.org/10.1016/B978-0-444-51624-4.50008-3>. Available at <http://www.sciencedirect.com/science/article/pii/B9780444516244500083>.

¹⁴ We do not call MBNF “abstract syntax,” because some of it is concrete syntax.

- [11] Donald E. Knuth (1986): *The TeXbook*. Addison-Wesley Professional.
- [12] Luis Fdo. Llana Díaz & Manuel Núñez (1997): *Testing semantics for unbounded nondeterminism*. In Christian Lengauer, Martin Griebel & Sergei Gorlatch, editors: *Euro-Par'97 Parallel Processing*, Springer Berlin Heidelberg, Berlin, Heidelberg, pp. 538–545.
- [13] Yiannis Moschovakis (1994): *Notes on Set Theory*, 1 edition. 0172-6056 978-1-4757-4153-7, Springer-Verlag New York, doi:10.1007/978-1-4757-4153-7.
- [14] John von Neumann (1923): *Zur Einführung der transfiniten Zahlen*. *Acta Scientiarum Mathematicarum* (Szeged) 1(4), pp. 199–208. Available at <http://acta.fyx.hu/acta/showCustomerArticle.action?id=4981&dataObjectType=article>. Auf Englisch nachgedruckt in GlossarWiki:Heijenoort:2002.
- [15] Benjamin C. Pierce (2002): *Types and Programming Languages*, 1st edition. The MIT Press.
- [16] V. Rahli, M. Bickford & R. L. Constable (2017): *Bar induction: The good, the bad, and the ugly*. In: *2017 32nd Annual ACM/IEEE Symposium on Logic in Computer Science (LICS)*, pp. 1–12, doi:10.1109/LICS.2017.8005074.
- [17] John C. Reynolds (2009): *Theories of Programming Languages*, 1st edition. Cambridge University Press, New York, NY, USA.
- [18] Helmut Seidl, editor (2012): *Programming Languages and Systems*. Springer.
- [19] Peter Sewell, Francesco Zappa Nardelli, Scott Owens, Gilles Peskine, Thomas Ridge, Susmit Sarkar & Rok Strniša (2007): *Ott: Effective Tool Support for the Working Semanticist*. *SIGPLAN Not.* 42(9), pp. 1–12, doi:10.1145/1291220.1291155. Available at <http://doi.acm.org/10.1145/1291220.1291155>.
- [20] Guy L. Steele, Jr. (2017): *It's Time for a New Old Language*. In: *Proceedings of the 22Nd ACM SIGPLAN Symposium on Principles and Practice of Parallel Programming, PPoPP '17*, ACM, New York, NY, USA, pp. 1–1, doi:10.1145/3018743.3018773. Available at <http://doi.acm.org/10.1145/3018743.3018773>.
- [21] Alfred Tarski (1955): *A lattice-theoretical fixpoint theorem and its applications*. *Pacific J. Math.* 5(2), pp. 285–309. Available at <https://projecteuclid.org:443/euclid.pjm/1103044538>.
- [22] Kazunori Tobisawa (2015): *A Meta Lambda Calculus with Cross-Level Computation*. In: *POPL '15*, pp. 383–393.
- [23] Neil Toronto & Jay McCarthy (2012): *Computing in Cantor's Paradise with λ ZFC*. In Tom Schrijvers & Peter Thiemann, editors: *Functional and Logic Programming*, Springer Berlin Heidelberg, Berlin, Heidelberg, pp. 290–306.
- [24] Norbert Wiener (1914): *A Simplification of the Logic of Relations*. *Proceedings of Cambridge Philosophical Society* 17, pp. 387–390. Nachgedruckt in GlossarWiki:Heijenoort:2002.

A Basic Logic and Mathematics

This appendix gives a brief overview of some concepts which are common enough in mathematics, but which are often represented in different ways, to say how they are used in this paper.

A.1 Metavariable Conventions

For this section, ν stands for an arbitrary metavariable (a meta-metavariable). Statements of the form “let ν range over C ” declare and define ν as a metavariable that stands for some element of the class C .

We use single letters (either Roman or Greek) for metavariables.

Whenever we declare a metavariable ν as ranging over a class, this also defines as ranging over that class all variants of ν obtained by either (1) adding a subscript $i \in \mathbb{N}$ to ν to produce ν_i (e.g., ν_0, ν_1, ν_2 ,

etc), (2) adding a single, double, or triple prime to ν , producing respectively in ν' , ν'' , and ν''' , or (3) a combination of (1) and (2).

In contrast, we use superscripts (e.g., ν^1 , ν^2) and accents (e.g., $\bar{\nu}$, $\tilde{\nu}$) to distinguish metavariables that are in some way related to the corresponding undecorated metavariable, but not necessarily ranging over the same class. For example, if we have declared ν to range over the set $\mathcal{S}$, we might have ν^0 ranging over $\mathcal{S}^0$, ν^1 ranging over $\mathcal{S}^1$, and $\mathcal{S}^1 \subset \mathcal{S}^0 \subset \mathcal{S}$.

A.2 Sets

The mathematical foundation we use is set theory with choice. ZFC is suitable, so are other variants. If $P(X)$ is a proposition of first-order logic that mentions X , then (1) $P(Y)$ differs from $P(X)$ only by mentioning Y instead of X , and (2) the notation $\{X \mid P(X)\}$ stands for $\{X \in \mathcal{S} \mid P(X)\}$ for some set $\mathcal{S}$ which is left to the reader to infer from the context of discussion. Given some expression $f(X_1, \dots, X_n)$ mentioning variables $X_1, \dots, X_n$, we use the notation $\{f(X_1, \dots, X_n) \mid P(X_1, \dots, X_n)\}$ for $\{Y \mid \exists X_1, \dots, X_n. Y = f(X_1, \dots, X_n) \wedge P(X_1, \dots, X_n)\}$. Given two sets X and Y we use the notation $X \perp Y$ to mean ‘ X and Y are disjoint.’

A.3 Pairs

We rely on an operator $(\cdot, \cdot)$ for building *ordered pairs* and corresponding *projection* operators fst and snd , such that if $Z = (X, Y)$, then $\text{fst}(Z) = X$ and $\text{snd}(Z) = Y$. We require that it is impossible for a pair to also be a set of pairs and that the natural numbers do not overlap with pairs.¹⁵ Given two sets $\mathcal{S}$ and $\mathcal{T}$, the *product set* $\mathcal{S} \times \mathcal{T}$ is the set of pairs $\{(X, Y) \mid X \in \mathcal{S} \text{ and } Y \in \mathcal{T}\}$. Let *tuple* notation be defined so that $(X_1, X_2, X_3, \dots, X_n) = ((X_1, X_2, X_3, \dots, X_{n-1}), X_n)$.

A.4 Relations

Let $\mathcal{R}$ range over sets of pairs. The statement $(X, Y) \in \mathcal{R}$ can be written with three kinds of alternate notation: $\mathcal{R}(X, Y)$, and $X \mathcal{R} Y$, and $X \xrightarrow{\mathcal{R}} Y$.

A relation $\mathcal{R}$ is *reflexive w.r.t. $\mathcal{S}$* iff $\mathcal{R} \supseteq \{(X, X) \mid X \in \mathcal{S}\}$. As is common practice, if we mention that a relation is reflexive without saying what set $\mathcal{S}$ this is with respect to, this means we are leaving it to the reader to infer from the context of discussion which set $\mathcal{S}$ to use.

Let $\mathcal{R}^*$ be the reflexive and transitive closure of $\mathcal{R}$ and let $\mathcal{R}^=$ be the reflexive, symmetric, and transitive closure of $\mathcal{R}$; in both cases we use the above-mentioned convention that the reader must infer the set $\mathcal{S}$ w.r.t. which to take the reflexive closure. Let $X \xrightarrow{\mathcal{R}} Y$ mean $X \xrightarrow{\mathcal{R}^*} Y$, and let $X \xleftrightarrow{\mathcal{R}} Y$ mean $X \xrightarrow{\mathcal{R}^=} Y$.

A relation is an *equivalence* iff it is symmetric and transitive. Given an equivalence relation $\mathcal{R}$, let $[X]_{\mathcal{R}} = \{Y \mid (X, Y) \in \mathcal{R}\}$ be the *equivalence class of X w.r.t. $\mathcal{R}$* and let $[X]_{\mathcal{R}}$ be an equivalence class of $\mathcal{R}$.

A relation $\mathcal{R}$ is *terminating* iff there is no infinite sequence $X_1, X_2, \dots$ such that $X_1 \xrightarrow{\mathcal{R}} X_2 \xrightarrow{\mathcal{R}} \dots$. If $X \xrightarrow{\mathcal{R}} Y$, and there exists no Z such that $Y \xrightarrow{\mathcal{R}} Z$, then we call Y an *$\mathcal{R}$ -normal form of X* . If $\mathcal{R}$ is terminating, then it can be used for *induction*: If it can be shown that $\mathcal{R}$ is terminating and $\forall X \in \mathcal{S}. (\forall Y \in \mathcal{S}. X \xrightarrow{\mathcal{R}} Y \Rightarrow P(Y)) \Rightarrow P(X)$, then it follows that $\forall X \in \mathcal{S}. P(X)$.

¹⁵We therefore can not use Kuratowski’s encoding of pairs where $(X, Y) = \{\{X\}, \{X, Y\}\}$, because (for example) $\{(X, X)\} = \{\{\{X\}\}\} = (\{X\}, \{X\})$. Similarly, we can not use the “short” encoding where $(X, Y) = \{X, \{X, Y\}\}$ together with von Neumann’s encoding of natural numbers (actually of all ordinal numbers) where $0 = \emptyset$ and $i + 1 = i \cup \{i\}$ because $(0, 0) = \{0, \{0, 0\}\} = \{\emptyset, \{\emptyset, \emptyset\}\} = \{\emptyset, \{\emptyset\}\} = \{\emptyset\} \cup \{\{\emptyset\}\} = 1 \cup \{1\} = 2$. We can use Wiener’s encoding of pairs where $(X, Y) = \{\{\{X\}, \emptyset\}, \{\{Y\}\}\}$, because in this encoding a pair can not be a set of pairs, a set of sets of pairs, or a von Neumann ordinal number. We can also work in a set theory with a primitive pairing operator.

A relation is a *partial order* on $\mathcal{S}$ iff it is transitive and antisymmetric. A partial order is *strict* iff it is irreflexive. A non-strict partial order, $\leq$, is a *total order* on $\mathcal{S}$ iff for all $X, Y \in \mathcal{S}$ either $X \leq Y$ or $Y \leq X$. A strict partial order, $<$, is a *strict total order* on $\mathcal{S}$ iff for all $X, Y \in \mathcal{S}$ s.t. $X \neq Y$ either $X < Y$ or $Y < X$.

A.5 Functions

A *function* is a relation f such that for all X, Y , and Z , if $\{(X, Y), (X, Z)\} \subseteq f$ then $Y = Z$. Let $\mathcal{S} \rightarrow \mathcal{T} = \{f \mid f \subseteq \mathcal{S} \times \mathcal{T} \text{ and } f \text{ is a function}\}$. Let f be *from* $\mathcal{S}$ *to* $\mathcal{T}$ iff $f \in \mathcal{S} \rightarrow \mathcal{T}$. A function f is *injective* iff f^{-1} is a function. If $(X, Y) \in f$ for some Y , then $f(X)$ denotes Y , otherwise $f(X)$ is undefined. A function f is *total* on $\mathcal{S}$ iff $f(X)$ is defined for all $X \in \mathcal{S}$. Given a function f , let $f[X \mapsto Y] = (f \setminus \{Z \in f \mid \text{fst}(Z) = X\}) \cup \{(X, Y)\}$.

A *fixed point* of a function f is some x for which $f(X) = X$. If the set of fixed points of f has a greatest lower bound which is itself a fixed point, then we call this the *least fixed point* of f and if it has a least upper bound which is itself a fixed point, then we call this the *greatest fixed point* of f .

A function is *f order preserving* w.r.t a partial ordering *leq* if $f(X) \leq f(Y)$ iff $X \leq Y$.

A.6 Sequences

Given a set $\mathcal{S}$ which is not a relation (if $\mathcal{S}$ contains only pairs then instead the notation refers to the definition of $\mathcal{R}^*$ from section A.4, the reflexive and transitive closure of $\mathcal{R}$), let $\mathcal{S}^*$, the set of *finite sequences of elements in* $\mathcal{S}$, be the set of all finite functions f such that $\text{range}(f) \subseteq \mathcal{S}$, and $\text{domain}(f) \subseteq \mathbb{N}$, and $m < n \in \text{domain}(f)$ implies $m \in \text{domain}(f)$.

Convention A.1 (Metavariables over Sequences). If v is declared to range over $\mathcal{S}$, then $\vec{v}$ is automatically declared to range over $\mathcal{S}^*$. $\square$

The notation $[v_0, \dots, v_n]$ stands for the least-defined function $\vec{v}$ such that $\vec{v}(i) = v_i$ for all $i \in \{0, \dots, n\}$. For example, the singleton sequence $[v]$ containing v as its only element is $\{(0, v)\}$, and we have $[v_0, v_1, v_2] = \{(0, v_0), (1, v_1), (2, v_2)\}$. The *component of a sequence* $\vec{v}$ *at index* i is simply $\vec{v}(i)$. Note that the first component of a sequence is at index 0, and that the *empty sequence* $[]$ is merely the empty set. The *length* of a sequence $\vec{v}$ is the smallest $n \in \mathbb{N}$ which is larger than all elements of $\text{domain}(\vec{v})$. The *concatenation* of sequences $\vec{v}_1$ and $\vec{v}_2$ is $\vec{v}_1 \cdot \vec{v}_2 = \vec{v}_1 \cup \{(|\vec{v}_1| + i, v) \mid (i, v) \in \vec{v}_2\}$.

Note that $(\mathcal{S}^*, \cdot, [])$ forms a monoid, i.e., the following equalities hold:

$$[] \cdot \vec{v} = \vec{v} \quad \vec{v} \cdot [] = \vec{v} \quad (\vec{v}_1 \cdot \vec{v}_2) \cdot \vec{v}_3 = \vec{v}_1 \cdot (\vec{v}_2 \cdot \vec{v}_3)$$

B Proof of Key Results

Section 2 is sufficient for any reader who wants an outline of our definition in order to interpret those pieces of MBNF it defines. This appendix will be of interest either to those readers who are looking to extend our definition to define more uses of MBNF, or to those readers who want to reassure themselves that the sort of entities described in section 2.1 can always be thought to exist. While it is our intention that our definition should be easy to work with, a deeper working knowledge of set theory is assumed for this appendix than the rest of this document and everything apart from section 3 may be read without this section.

For this appendix we use Wiener's [24] encoding of pairs and von Neumann's encoding of ordinals, the natural numbers [14] and cardinal assignment [13]. We also make use of the axiom of choice.

Lemma B.1. *Given a countable set A and a set B of all trees C such that the interior nodes of C are elements of A and the leaf nodes of C are elements of ω_1 , $|B| = \aleph_1$.*

Proof. Let D be the set of all trees C such that every element of C is an element of ω_1 . Since we can make a bijection between members of A to a members of ω and the function $f \in \omega_1 \rightarrow \omega_1$ s.t. $f(x) = \omega + x$ is a bijection, $|C| = |D|$. For a finite subset, S , of ω_1 the relation $<$ on S is a finite subset of $\omega_1 \times \omega_1$. Assuming choice, $|\omega_1 \times \omega_1| = \aleph_1$. The cardinality of the set of finite subsets of ω_1 is $\aleph_1$. So $|C| = |D| = \aleph_1$. $\square$

We can now go on to show that Object and Arrangement are well defined. We do so by producing a model within set theory that fulfils most of the constraints in section 2.1., which are written out formally in Appendix D.¹⁶ This proof requires that the reader pick some appropriate values for Symbol, Pos, Pointer, $\square$, ϵ and $\bar{B}$. The definition of these sets in Appendix D is adequate.

We for a given ordinal i we define a set of tuples OPAC_i which may be thought of as getting closer to the tuple (Object, ptr, Arrangement, Core).

Definition B.2 (OPAC).

0 Case:

Let $\text{Obj}_0 = \{\square\}$

Let $\text{ptrSpace}_0 = \{x \in \text{Obj}_0 \rightarrow \text{Pointer} \mid x \text{ is total on } \text{Obj}_0 \wedge x \text{ is injective}\}.$

Let

$$\text{OPAC}_0 = \{(\text{Obj}_0, \text{ptr}_0, \text{Arr}_0, \text{Core}_0) \mid \begin{array}{l} \text{ptr}_0 \in \text{ptrSpace}_0 \wedge \text{Core}_0 = \text{Arr}_0 \setminus \{\epsilon\} \\ \wedge \text{Arr}_0 = \mathbb{N} \cup \{\epsilon\} \cup \text{Symbol} \cup \text{ptr}_0(x) \end{array}\}$$

+1 Case:

For $(\text{Obj}_n^k, \text{ptr}_n^k, \text{Arr}_n^k, \text{Core}_n^k) \in \text{OPAC}_n$

Let $\text{Accent}_{n+1}^k = (\{\bar{B}\} \times \text{Symbol}) \times (\text{Arr}_n^k \setminus \{\epsilon\}).$

Let $\text{Core}_{n+1}^k = \text{Core}_n^k \cup \text{Accent}_{n+1}^k.$

Let $\text{Layout}_{n+1}^k = \text{Arr}_n^k \times \{x \in \text{Pos} \rightarrow \text{Arr}_n^k \setminus \{\epsilon\} \mid x \neq \emptyset\}.$

Let $\text{Seq}_{n+1}^k = \text{Arr}_n^k \times \text{CoreArr}_{n+1}^k.$

Let $\text{Obj}_{n+1}^k = \text{Obj}_n^k \cup \{x \in \mathcal{P}(\text{Arr}_n^k) \mid |x| \leq \aleph_0\}$

Let

$$\text{ptrSpace}_{n+1} = \{x \in \text{Obj}_{n+1}^k \rightarrow \text{Pointer} \mid \begin{array}{l} (\text{Obj}_n, \text{ptr}_n, \text{Arr}_n, \text{Core}_n) \in \text{OPAC}_n \wedge \\ x \text{ is total on } \text{Obj}_{n+1}^k \wedge x \text{ is injective} \end{array}\}$$

Let $\text{ptr}_{n+1}^{k,i}(x) \in \text{ptrSpace}_{n+1}$ such that $\text{ptr}_{n+1}^{k,i}(x) \subseteq \text{ptr}_{n+1}^k$ if such a set exists. If no such set exists let $\text{ptr}_{n+1}^{k,0} = \emptyset$

Let $\text{Arr}_{n+1}^{k,i} = \text{Arr}_n^k \cup \text{CoreArr}_{n+1}^k \cup \text{Layout}_{n+1}^k \cup \text{Seq}_{n+1}^k \cup \{\text{ptr}_{n+1}^{k,i}(x) \mid x \in \text{Obj}_{n+1}^k\}$ if $\text{ptr}_{n+1}^{k,i}(x)$ is defined and $\emptyset$ otherwise.

Let

$$\text{OPAC}_{n+1} = \{(\text{Obj}_{n+1}^k, \text{ptr}_{n+1}^{k,i}, \text{Arr}_{n+1}^{k,i}, \text{Core}_{n+1}^k) \mid \begin{array}{l} (\text{Obj}_n, \text{ptr}_n, \text{Arr}_n, \text{Core}_n) \in \text{OPAC}_n \\ \wedge \text{Arr}_{n+1}^{k,i} \neq \emptyset \end{array}\}$$

Limit Case:

¹⁶In order to simplify the proof this model allows the use of arrangements consisting of a single pointer in the formation of objects. It is not too difficult to rule this case out.

We now define the above functions for a limit point ε .

Let

$$\begin{aligned} \text{stack} = \{S \subseteq \bigcup_{i < \varepsilon} \text{OPAC}_i \mid & ((\text{Obj}_i^k, \text{ptr}_i^k, \text{Arr}_i^k, \text{Core}_i^k) \in S \wedge (\text{Obj}_j^l, \text{ptr}_j^l, \text{Arr}_j^l, \text{Core}_j^l) \in S) \\ \Rightarrow & ((j < i \Rightarrow (\text{Obj}_j^l \subseteq \text{Obj}_i^k \wedge \text{Arr}_j^l \subseteq \text{Arr}_i^k \wedge \text{ptr}_j^l \subseteq \text{ptr}_i^k)) \\ & \wedge (j < i \vee i < j \\ & \vee (\text{Obj}_i^k, \text{ptr}_i^k, \text{Arr}_i^k, \text{Core}_i^k) = (\text{Obj}_j^l, \text{ptr}_j^l, \text{Arr}_j^l, \text{Core}_j^l)) \\ & \wedge \forall n < \varepsilon, (\text{Obj}_n^m, \text{ptr}_n^m, \text{Arr}_n^m, \text{Core}_n^m) \in S\} \end{aligned}$$

For $S \in \text{stack}$

Let

$$\text{Obj}_\varepsilon^S = \left(\bigcup \{ \text{Obj}_i \mid (\text{Obj}_i^k, x, y, z) \in S \} \right) \cup \left\{ x \in \mathcal{P} \left(\bigcup \{ \text{Arr}_i^k \mid (a, b, \text{Arr}_i^k, c) \in S \} \right) \mid |x| \leq \aleph_0 \right\}$$

Let $\text{ptr}_\varepsilon^{S,k}$ be a bijection between $S_\varepsilon \subseteq \text{Pointer}$ and Obj_ε^S such that for all $(\text{Obj}_i, \text{ptr}_i, \text{Arr}_i, \text{Core}_i) \in S$, $\text{ptr}_i \subseteq \text{ptr}_\varepsilon^{S,k}$. If no such bijection exists, let $\text{ptr}_\varepsilon^{S,0} = \emptyset$.

Let $\text{Arr}_\varepsilon^{S,k} = \{ \text{ptr}_\varepsilon^{S,k}(x) \mid x \in \text{Obj}_\varepsilon^S \} \cup \{ \text{Arr}_i^k \mid (a, b, \text{Arr}_i^k, c) \in S \}$ if $\text{ptr}_\varepsilon^{S,k}(x)$ is defined and $\emptyset$ otherwise.

Let $\text{Core}_\varepsilon^S = \bigcup \{ \text{CoreArr}_i \mid (x, y, z, \text{Core}_i^k) \in S \}$.

Let

$$\text{OPAC}_\varepsilon = \left\{ (\text{Obj}_\varepsilon^S, \text{ptr}_\varepsilon^{S,k}, \text{Arr}_\varepsilon^{S,k}, \text{Core}_\varepsilon^S) \mid S \in \text{stack} \wedge \text{Arr}_{n+1}^{k,i} \neq \emptyset \right\}$$

□

Lemma B.3. *OPAC_i is Non-Empty for all Ordinals i.*

Proof. The only way OPAC_i may be empty for some i is if $|\text{Obj}_i^k| > \text{Pointer}$ for some Obj_i^k such that $(\text{Obj}_i^k, a, b, c) \in \text{OPAC}_i$. We prove by induction on the size of Obj_n^k and Arr_n^k such that $(\text{Obj}_n^k, x, \text{Arr}_n^k, y) \in \text{OPAC}_n$ that this cannot be the case.

0 Case:

$|\text{Obj}_0| = 1 \leq \aleph_1$ and for all $\text{Arr}_0 \in \text{ArrSpace}_0$, $|\text{Arr}_0| = \aleph_0 \leq \aleph_1$.

+1 Case:

If, for all $(\text{Obj}_n^k, x, \text{Arr}_n^k, y) \in \text{OPAC}_n$, $|\text{Obj}_n^k| \leq \aleph_1$ and $|\text{Arr}_n^k| \leq \aleph_1$, then, for all Obj_{n+1}^k , $|\text{Obj}_{n+1}^k| \leq \aleph_1$, provided we have some way of ordering Obj_i^k and Arr_i^k . With choice this follows quite easily from the fact that the cardinality of the set of subsets of $\aleph_1$ which are of cardinality less than or equal to $\aleph_0$ is $(2^{\aleph_0})^{\aleph_0} = 2^{\aleph_0 \cdot \aleph_0} = 2^{\aleph_0}$. As, for all Obj_{n+1}^k there exists some Obj_n^j s.t. $\text{Obj}_{n+1}^k \subseteq \text{Obj}_n^j$, there exists some ptr which assigns pointers for Obj_{n+1}^k and which may also be used to assign pointers for Obj_n^j .

It is easy to observe that, if $|\text{Obj}_n^k| \leq \aleph_1$ and $|\text{Arr}_n^k| \leq \aleph_1$, then $|\text{Arr}_{n+1}^{k,i}| \leq \aleph_1$ since neither Accent_{n+1}^k , nor Layout_{n+1}^k nor Seq_{n+1}^k can add cardinality greater than $\aleph_1$.

Limit Case:

We show $\exists \varepsilon; \forall i < \varepsilon; (|\text{Obj}_i| \leq \aleph_1 \wedge |\text{Arr}_i| \leq \aleph_1) \Rightarrow (|\text{Obj}_\varepsilon| \leq \aleph_1 \wedge |\text{Arr}_\varepsilon| \leq \aleph_1)$. We note that no Arr_i has $\aleph_0$ sub-arrangements and all such Arr_i can be readily identified with some finite tree whose interior nodes are labelled corresponding the operations accenting, concatenation and the finite number of possible combinations of Subscript, superscript etc. and whose leaf nodes are labelled with members of the set ω_1 . So, by lemma B.1, $|\bigcup_{i=0}^{\varepsilon} \text{Arr}_i| \leq \aleph_1$. Similarly we may readily identify each set in Obj_i apart from the $\square$ with some countable subset of the set of trees we used to define each Arr_i . The cardinality of the countable subsets of a set of size $\aleph_1$ is $(2^{\aleph_0})^{\aleph_0} = 2^{\aleph_0 \cdot \aleph_0} = 2^{\aleph_0}$. So $|\bigcup_{i=0}^{\varepsilon} \text{Obj}_i| \leq \aleph_1$. The desired result

follows easily. As $\text{Obj}_\varepsilon \subseteq \text{Obj}_i$, for each $i \leq \varepsilon$ there exists some ptr which assigns pointers for Obj_ε and which may also be used to assign pointers for Obj_i . $\square$

Definition B.4 (fun). Let $Z = \{\text{OPAC}_i \mid i < \kappa\}$ for some $\kappa < \omega_2$. We define a function $\text{fun} \in Z \rightarrow Z$ such that $\text{fun}(\text{OPAC}_i) = \text{OPAC}_{i+1}$ $\square$

Lemma B.5. *fun has a least fixed point.*

Proof. The Knaster–Tarski theorem [21] tells us that any order-preserving function on a complete lattice has a least fixed point. For $\text{OPAC}_a, \text{OPAC}_b \in Z$ we define $\leq$ such that $\text{OPAC}_a \leq \text{OPAC}_b$ iff (either, for all $(\text{Obj}_b, p, q, r) \in \text{OPAC}_b$, there exists $(\text{Obj}_a, b, c, d) \in \text{OPAC}_a$ s.t. $\text{Obj}_a \subset \text{Obj}_b$, or, for all $(\text{Obj}_b, p, \text{Arr}_b, r) \in \text{OPAC}_b$, there exists $(\text{Obj}_a, b, \text{Arr}_a, d) \in \text{OPAC}_a$ s.t. $(\text{Obj}_a = \text{Obj}_b$ and $\text{Arr}_a \subset \text{Arr}_b)$). Z is a complete lattice ordered by $\leq$ and fun is an order preserving function on Z .¹⁷ $\square$

Theorem B.6. *Object and Arrangement are well defined.*

Proof. For some tuple a in OPAC_i there exists some tuple b in OPAC_i such that:

1. The first member of b contains all the objects our rules say must exist if Arrangement is at least the third member of a ,
2. The third member of b contains all the arrangements that our rules say must exist if Arrangement is at least the third member of a , Object is at least the first member of a and ptr is at least the second member of a .

If OPAC_i is the least fixed point of fun then $\text{OPAC}_{i+1} = \text{OPAC}_i$.

We now take the least fixed point, $\text{lfp}(\text{fun})$, of $\text{fun} \in Z \rightarrow Z$ and select some tuple $(\text{Obj}, \text{ptr}, \text{Arr}, \text{Core}) \in \text{lfp}(\text{fun})$. The first member of the tuple gives us a model for Object and the third member, Arrangement. $\square$

C Examples of our Definition in Action

C.1 Call by Need

The following example is derived from Chang and Felleisen [2, p 134]:

$$\begin{array}{ll}
 e \in se ::= x \mid \lambda x. e \mid e e & \hat{A} \in s\hat{A} ::= \square \mid A[\hat{A}] e \\
 v \in sv ::= \lambda x. e & \check{A} \in s\check{A} ::= \square \mid A[\lambda x. \check{A}] e \\
 a \in sa ::= A[v] & E \in sE ::= \square \mid E e \mid A[E] \mid \hat{A}[A[\lambda x. \check{A}[E[x]]]E] \\
 A \in sA ::= \square \mid A[\lambda x. A] e & \text{where } \hat{A}[\check{A}] \in A
 \end{array}$$

Each constraint is added sequentially and the least set of objects satisfying them is recalculated. Where the value of a set a metavariable can range over is recalculated and it is referenced in another rule, the set that rule applies to is recalculated with a new value. For example, initially $a \in sa = \emptyset$, but when $A ::= \square$ is read it triggers a recalculation of $A[v]$ so $sa = sv$. Then when $A ::= A[\lambda x. A] e$ is read, first it triggers a recalculation of A so $sA = \{\square\} \cup \{[P_{\lambda A} P_e]_{\approx} \mid \text{ptr}(e) = P_e \wedge \text{ptr}([\lambda x. \square]_{\approx}) = P_{\lambda A}\}$ then it triggers a recalculation of a so $sa = v \cup \{[P_{\lambda A[v]} P_e]_{\approx} \mid \text{ptr}(e) = P_e \wedge \text{ptr}([\lambda x. P_v]_{\approx}) = P_{\lambda A[v]} \wedge \text{ptr}(v) = P_v\}$ then

¹⁷Note that the way in which we have defined Obj and Arr is such that $j \leq i$ implies $(\text{Obj}_j \subseteq \text{Obj}_i$ and $\text{Obj}_j \subseteq \text{Obj}_i)$. Note also that, for all limit ordinals $\varepsilon \leq \kappa$; $(\text{Obj}_\varepsilon, \text{Arr}_\varepsilon) \in Z$. Finally note that ω_2 is large enough that it has a larger cardinality than any Obj_i or Arr_i , so we can select some κ larger than the partition of Obj_i and Arr_i into the extra elements added at each stage.

a won't trigger recalculations, but we have a recalculation on A waiting. Let $f(se, sv, sa, sA, s\hat{A}, s\check{A}, sE)$ take $(se, sv, sa, sA, s\hat{A}, s\check{A}, sE)$ when a recalculation is triggered to their values after a recalculation is performed. Let $<$ be an relation on $(se, sv, sa, sA, s\hat{A}, s\check{A}, sE)$ such that $(se^1, sv^1, sa^1, sA^1, s\hat{A}^1, s\check{A}^1, sE^1) < (se^2, sv^2, sa^2, sA^2, s\hat{A}^2, s\check{A}^2, sE^2)$ iff $se^1 \subset se^2$ or $sv^1 \subset sv^2$ or $sa^1 \subset sa^2$ or $sA^1 \subset sA^2$ or $s\hat{A}^1 \subset s\hat{A}^2$ or $s\check{A}^1 \subset s\check{A}^2$ or $sE^1 \subset sE^2$. We observe that each set out of $(se, sv, sa, sA, s\hat{A}, s\check{A}, sE)$ either gets new elements added to it or remains the same every time a recalculation is triggered and is bounded above by **Object**. We can therefore take the least fixed point on f satisfying all of these constraints.

If the side condition on E were to re-trigger the calculation on A , $\hat{A}$ or $\check{A}$ we would have to be able to check that the new side condition produced by this recalculation could not effect any E previously added. This grammar relies on an assignation of values to x , otherwise all of its sets are either $\emptyset$ or $\square$. For a given equivalence (e.g. $\equiv_\alpha$ or $\equiv_A$) this grammar may define a different collection of $(se, sv, sa, sA, s\hat{A}, s\check{A}, sE)$.

We add the reduction rule for this Grammar which is the least $\mathcal{R}$ satisfying:

$$\hat{A}[A_1[\lambda x. \check{A}[E[x]]]A_2[v]] \xrightarrow{\mathcal{R}} (\hat{A}[A_1[A_2[(\check{A}[E[x]])[x := v]]]]) \quad \text{where } \hat{A}[\check{A}] \in A$$