

Structured Document Algebra in Action

Don Batory¹, Peter Höfner², Dominik Köppl³,
Bernhard Möller⁴, and Andreas Zelend⁴

¹ Dept. of Computer Science, University of Texas at Austin, USA

² NICTA and UNSW, Australia

³ Department of Computer Science, TU Dortmund, Germany

⁴ Institut für Informatik, Universität Augsburg, Germany

Abstract. A *Structured Document Algebra (SDA)* defines modules with variation points and how such modules compose. The basic operations are module addition and replacement. Repeated addition can create nested module structures. SDA also allows the decomposition of modules into smaller parts. In this paper we show how SDA modules can be used to deal algebraically with *Software Product Lines (SPLs)*. In particular, we treat some fundamental concepts of SPLs, such as refinement and refactoring. This leads to mathematically precise formalization of fundamental concepts used in SPLs, which can be used for improved *Feature-Oriented Software Development (FOSD)* tooling.

Keywords: software product lines, feature-oriented design, algebraic reasoning.

It is our pleasure to dedicate this paper to Martin Wirsing on the occasion of his Formal Retirement. We contribute a study on a recently developed algebra for all kinds of structured and interconnected documents, but particularly the ones that describe product families or product lines in Feature Oriented Software Design — a topic on which Martin has been quite active for a while now. The frame of this are general formal methods and semantics. We pick up this latter theme to endow the algebra, which previously had a more syntactic flavor, with a semantic component, too. The particular approach we take uses the terminology of algebraic specification; this gives the fourth author the opportunity to fondly remember the days of the CIP project at TU Munich, in which Martin and he cooperated quite a lot on that topic. We hope that Martin will enjoy that bit of scientific nostalgia, too! Best wishes, Martin, for your Formal Retirement — enjoy! — but also for many further successful years, since we do hope that your Retirement is only Formal!

1 Introduction

A *Software Product Line (SPL)* is a family of related programs constructed from a common set of assets. Variations in programs are explained by *features*—increments in program functionality. The assets of an SPL are modules that implement features. These modules are the building blocks of SPL programs [2].

Today’s SPL researchers are exploring two rather different forms of feature-based modularity: alternative-based variation (a.k.a. *classical modularity*) and projectional variation (a.k.a. *SYSGEN* or *virtual modularity*). Classical modularity is what you would expect: there are physical files that define a feature module and tools that compose modules to produce a desired program. Here, files refer to arbitrary documents, such as specifications, code composed of elementary program features, text fragments or manuals. In contrast, virtual feature modularity is a preprocessor technology called *coloring*. The idea is simple: the code of the **Blue** feature is painted blue; code of the **Green** feature is painted green. Whenever **Blue** is not needed in a product, all blue-colored code is removed or is said to be *projected out*. The tools for virtual modularity are historically based on text preprocessors; more advanced tools color *abstract syntax trees (ASTs)* [5,15,18]. The current debate is which implementation technique is most appropriate for an SPL. Our position is that both implement the same abstractions—feature modules—in very different ways. *What is important is to understand the algebraic nature of these abstractions.*

The *Structured Document Algebra (SDA)*, partially first presented in [6], aims at providing a simple, yet effective, algebra of feature modularity (e.g., the modularization of text files, text fragments, ASTs, etc.).

It is completely independent of the underlying programming paradigm, such as Object-Oriented Programming or Functional Programming. In fact, it is even independent of programming, since it is abstract enough to cover also general structured documents, such as manuals or collections of web pages. It is meant as an aid for formal reasoning about the process of decomposing larger document pieces into smaller ones (and vice versa). It can be implemented on top of or inside any text editor, IDE or web page editor. Depending on the desired level of reasoning, it can be used purely syntactically or at a semantic level, as demonstrated in the Appendix. Phenomena modelled by the algebra include coherence, uniform transformations, deletion, overriding. In the special case of SPL documents, SDA additionally allows a description of their fine structure. This is in contrast to other algebraic approaches in this area (e.g., [1,24,25]), where modules are often treated as atomic units.

In the present paper we extend the basic repertoire of the version of SDA from [6] and show how it can be used to formally describe some standard techniques used in feature oriented software construction and SPLs.

The basis of SDA are structured, inter-linked documents or document *fragments*. A link is represented as a *Variation Point (VP)*, i.e., a labeled position in a fragment where contents can be inserted to yield a larger fragment. The association between VPs and their assigned fragments, if any, is provided by *modules*, i.e., partial functions from VPs to fragments.

SDA allows multiple “applied occurrences”, i.e., replication of one and the same VP v ; they all stand for (or share) the fragment assigned to v by some module. Conversely, different VPs may be associated with the same fragment (a module need not be an injective partial function).

In this paper we show how modules and fragments can be treated algebraically and discuss algebraic operations for module (de)composition and present an operation for overriding.

2 Structured Document Algebra

This section recapitulates a formal model of VPs, modules containing VPs, and compositions of such modules first presented in [6]. To keep SDA language-independent, we leave the exact nature of fragments (e.g., text or abstract syntax trees) unspecified and view it as a parameter of the algebra. For our examples we will use Java code fragments.

2.1 SDA Basics

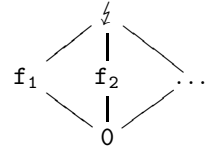
Variation Points and Fragments. We now formalize the notions mentioned above. Let V be a set of VPs, denoted by $v_1, v_2, \dots$ at which fragments may be inserted and $F(V)$ be a set of *fragments*, denoted by $f_1, f_2, \dots$. Fragments may contain VPs from V .

In incremental software design it is often advantageous to leave certain parts unspecified and to insert placeholders where (optionally) further features may be added. As an abstraction of such placeholders we use *default fragments*. In addition, it is convenient to introduce a special pseudo-fragment that represents an error, namely that there has been an attempt to assign two or more non-default and different fragments to the same VP. To this end we assume that the set $F(V)$ includes two special elements, a default fragment 0 and an error ζ .

The addition, or supremum, operator $+$ on fragments has the axioms

$$0 + x = x, \quad \zeta + x = \zeta, \quad f_i + f_j = \zeta \quad (i \neq j),$$

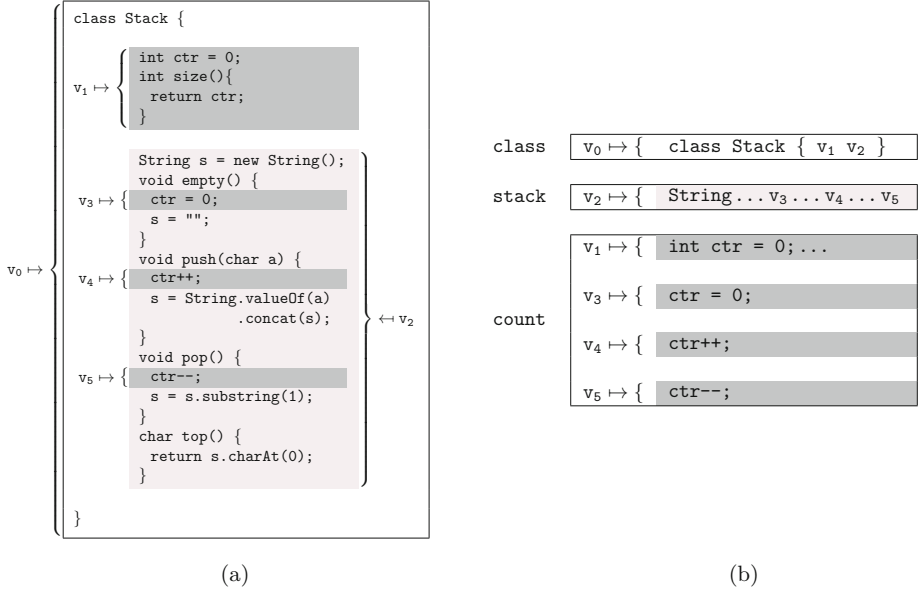
where $x, f_i, f_j \in F(V)$ with $f_i, f_j \neq 0$. If we assume associativity, idempotence and commutativity of addition, this structure forms a flat lattice with least element 0 and greatest element ζ .



Modules. A *module* is a partial function $m : V \rightsquigarrow F(V)$. A VP v is *assigned* by m if $v \in \text{dom}(m)$, otherwise *unassigned* or *external*. Thus the domain $\text{dom}(m)$ of a module is the set of VPs it “knows about” or administers. By using partial functions rather than relations, a VP can be filled with at most one fragment (*uniqueness*).

Example 2.1. Figure 1(a) is a sample file/module, structured by the assignment of fragments to its VPs. Its partial function is given in Figure 1(b). $\square$

Ducasse et al. [12] also use a flat lattice of composable units, called *traits*. Our modules correspond to the *method dictionaries* there. However, these have to be total rather than partial maps, which makes distinguishing assigned and external VPs difficult. Had we taken the same decision, our algebraic laws would have become much more cumbersome.

**Fig. 1.** Variation Points, Fragments and Modules

A module m can be viewed in different ways:

- as a collection of fragments that instantiate the VPs of $\text{dom}(m)$, i.e., a structured document;
- as filling certain VPs with contents (in term rewriting etc., it would be called a *substitution*); and
- as a generalized context-free grammar with $\text{dom}(m)$ as the set of nonterminals and a production $v \mapsto m(v)$ for each $v \in \text{dom}(m)$.

The simplest module is the *empty module* 0 , i.e., the empty partial map. Another very simple kind of module is provided by *constant modules*, i.e., modules which assign one and the same fragment—for instance 0 —to a number of VPs. Let $W \subseteq V$ be a set of VPs and f a fragment. We set

$$[W \mapsto f](v) =_{df} \begin{cases} f & \text{if } v \in W, \\ \text{undefined} & \text{otherwise.} \end{cases}$$

If $W = \{w\}$ is a singleton set, we abbreviate $[\{w\} \mapsto f]$ to $[w \mapsto f]$. Such *singleton modules* are the atomic building blocks from which all other modules can be constructed by the addition operation introduced next.

Module Addition. We want to construct larger modules step by step by assigning more and more fragments to VPs. The central operation for this is module addition ($+$), which fuses two modules while maintaining uniqueness (and signaling an error upon a conflict). Basically, module addition can be viewed as union with flagged inconsistent VP/fragment associations. Desirable properties for $+$ are

commutativity and associativity. If two modules have no VPs in common, the partial functions characterizing them can easily be combined. For example, `class + stack` (Figure 1(b)) is the partial function

$$\text{class} + \text{stack} = \begin{cases} v_0 \mapsto \text{class Stack } \{v_1 \ v_2\} \\ v_2 \mapsto \text{String } \dots v_3 \dots v_4 \dots v_5 \end{cases}$$

Module addition can be defined as the lifting of $+$ on fragments

$$(\mathbf{m} + \mathbf{n})(v) =_{df} \begin{cases} \mathbf{m}(v) & \text{if } v \in \text{dom}(\mathbf{m}) - \text{dom}(\mathbf{n}) \\ \mathbf{n}(v) & \text{if } v \in \text{dom}(\mathbf{n}) - \text{dom}(\mathbf{m}) \\ \mathbf{m}(v) + \mathbf{n}(v) & \text{if } v \in \text{dom}(\mathbf{m}) \cap \text{dom}(\mathbf{n}) \\ \text{undefined} & \text{if } v \notin \text{dom}(\mathbf{m}) \cup \text{dom}(\mathbf{n}) \end{cases}$$

If in the third case $\mathbf{m}(v) \neq \mathbf{n}(v)$ and $\mathbf{m}(v), \mathbf{n}(v) \neq 0$ then $(\mathbf{m} + \mathbf{n})(v) = \perp$, thus signaling an error. By the above laws for fragment addition, the set of modules forms a commutative and idempotent monoid under $+$.¹ Therefore we can define a *submodule relation* as the natural order

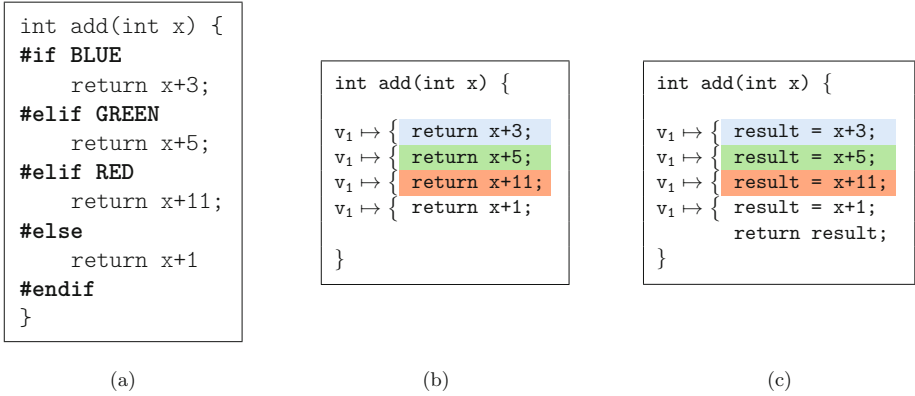
$$\mathbf{m} \leq \mathbf{n} \Leftrightarrow_{df} \mathbf{m} + \mathbf{n} = \mathbf{n} . \quad (1)$$

Example 2.2. Figure 1(b) shows three modules: `class` and `stack` contain single fragments, assigned to v_0 and v_2 respectively; `count` contains fragments assigned to v_1 , v_3 , v_4 , and v_5 . The module addition `class + stack + count` corresponds to the code in Figure 1(a). $\square$

Implementation. A simple example suggests several ways in which SDA modules can be implemented. Figure 2(a) shows how preprocessor macros might define three non-default fragments (labeled BLUE, GREEN, RED) and a default for an implicit variation point. Figure 2(b) shows how this might be rendered in a “coloring” tool (e.g., the one contained in the integrated development environment CIDE [18,19,2]) where the fragments assigned to the VPs are explicitly shown. (There is no need to actually “see” the names of VPs). However, Figure 2(b) would require significant engineering: a Java compiler would have to understand the preprocessor semantics of coloring (Figure 2(a)) so as not to alert programmers that the GREEN fragments and beyond are unreachable if BLUE is true. In a Java-like language this can be accomplished by defining feature variables as static Booleans and put a wrapper `if (feature){ fragment }` around each fragment. So if the feature variable is false, the compiler will effectively ignore the fragment as dead code. This is, for instance, offered by the CIDE compiler.

A more likely possibility—which is consistent with current text coloring tools—would be to “fool” the compiler by pretending that the code of Figure 2(c) is the definition of the `add` method, where a projection would produce a simpler method with only one assignment to the variable `result`.

¹ Modules and module addition can be recoded in terms of total functions, which makes it easier to see that the $+$ operation indeed is commutative, associative and idempotent, hence induces a lattice, too. Moreover, it has the empty module 0 as its neutral element and satisfies $\text{dom}(\mathbf{m} + \mathbf{n}) = \text{dom}(\mathbf{m}) \cup \text{dom}(\mathbf{n})$.

**Fig. 2.** Module Implementations

These ideas are, in effect, standard fare for SPL development, except that the tool support needs to be beautified by coloring and VP recognition. Code fragments or *mini-modules* can indeed be expressed in terms of classical module systems; see [10] for examples. Coloring also presents the connection between modules and virtual modularity (cf. Sect. 1).

2.2 Structural Properties of Modules

Since fragments may contain VPs, cycles could occur when composing modules. A module, however, should be cycle-free. To handle this we use a dependence relation.

Cycle-Freeness. For a fragment $f \in F(V)$ let $VP(f)$ be the set of VPs that occur in f . We define a *direct dependence relation* $\text{dep}_m \subseteq V \times V$ within a module m by

$$v \text{ dep}_m w \Leftrightarrow_{df} v \in \text{dom}(m) \wedge w \in VP(m(v)) .$$

This means that VP w occurs in the fragment assigned to VP v by m , so that v depends on w . For example, in Figure 1(b), $v_2 \text{ dep}_{\text{stack}} v_3$. A module m is *acyclic* if no VP depends directly or indirectly on itself, i.e., no VP v satisfies $v \text{ dep}_m^+ v$, where dep_m^+ is the transitive closure of dep_m . Henceforth we only consider cycle-free modules.

The concepts of dependence and module addition do not interfere. Since the latter is just the “union” of VP/fragment associations, for acyclicity it does not matter whether dependences occur inside one module or between different modules.

Assembling Fragments. We now describe how to assemble a structured document into a single fragment (while “forgetting” the structure). To define this formally we use an auxiliary function `single_fill(f,m)`. It takes a fragment f and a module m and yields the fragment that results from f by replacing, in parallel, all occurrences of every $v \in VP(f)$ by the corresponding fragment $m(v)$

(if any). The precise definition of `single_fill` depends on the special type of fragments considered; as stated in the introduction we want to keep that parametric. For an acyclic module m and $v \in \text{dom}(m)$, the fragment $\text{frag}(v, m)$ can be computed by iterating the `single_fill` function. By the assumed acyclicity of m this always terminates. To cope with unassigned (= external) VPs we assume that every VP is also a fragment, i.e., that $v \subseteq F(v)$. Then the external VPs can simply be left unchanged by the assembly function. A corresponding program looks as follows:

```
fragment frag (vp v,module m){
    fragment f = v;
    while (VP(f) ∩ dom(m) ≠ ∅)
        f = single_fill(f,m);
    return f; }
```

Note that $\text{frag}(v, m) = v$ if $v \notin \text{dom}(m)$.

Module Equivalence. Two modules m, n are *equivalent* if they generate the same fragments for all VPs:

$$m \equiv n \Leftrightarrow_{df} \forall v \in VP : \text{frag}(v, m) = \text{frag}(v, n) .$$

Lemma 2.3. *For modules m, n we have $m \equiv n \Rightarrow \text{dom}(m) = \text{dom}(n)$.*

Restructuring. Module m *restructures into* module n , in symbols $m \sqsubseteq n$, if $m \sqsubseteq n \Leftrightarrow_{df} \text{dom}(m) \subseteq \text{dom}(n) \wedge \forall v \in \text{dom}(m) : \text{frag}(v, m) = \text{frag}(v, n)$.

This means that n offers a possibly more detailed representation of the fragments of m using auxiliary VPs in $\text{dom}(n) - \text{dom}(m)$. The empty module is the smallest one w.r.t. $\sqsubseteq$, since it does not offer any choice—in particular since $\text{dom}(\mathbf{0}) = \emptyset$.

It is easy to check that the relation $\sqsubseteq$ is reflexive and transitive, i.e., a preorder. It is, however, not antisymmetric as the following example shows.

Example 2.4. Consider the modules m and n defined as follows:

$$\begin{array}{ll} m : & v \mapsto i++; w; \quad \text{and} \quad w \mapsto j++; \\ n : & v \mapsto i++; j++; \quad \text{and} \quad w \mapsto j++; \end{array}$$

These are equivalent, i.e., $m \equiv n$, in that they produce the same text. Hence, $m \sqsubseteq n$ and $n \sqsubseteq m$. If $\sqsubseteq$ were antisymmetric, this should imply $m = n$, but m and n are different as modules. $\square$

However, we always have $m \sqsubseteq n \wedge n \sqsubseteq m \Leftrightarrow m \equiv n$.

We provide some basic properties of $\sqsubseteq$ that can be used to restructure larger modules in a modular fashion. To state them we define the set of all VPs occurring in a module m by $VP(m) =_{df} \text{dom}(m) \cup \bigcup_{v \in \text{dom}(m)} VP(m(v))$. First we treat the case of decomposing the fragment belonging to a single VP.

Lemma 2.5. *Assume fragments f, g and a module m with $\text{dom}(m) = VP(g)$ and $f = \text{single_fill}(g, m)$. This means that the VPs in g are filled by m yielding f . Moreover let v be a fresh VP, i.e., $v \notin \text{dom}(m) \cup VP(m)$. Then*

$$[v \mapsto f] \sqsubseteq ([v \mapsto g] + m) .$$

The proof is immediate from the definitions and assumptions.

Lemma 2.6. *For modules $\mathbf{m}_i \sqsubseteq \mathbf{n}_i$ ($i = 1, 2$) with $\text{VP}(\mathbf{n}_1) \cap \text{VP}(\mathbf{n}_2) = \emptyset$ we have*

$$\mathbf{m}_1 + \mathbf{m}_2 \sqsubseteq \mathbf{n}_1 + \mathbf{n}_2 .$$

Proof. (Sketch) We first state an auxiliary property. Assume modules $\mathbf{m}, \mathbf{n}$ such that $\mathbf{m} + \mathbf{n}$ is acyclic and $\text{VP}(\mathbf{m}) \cap \text{VP}(\mathbf{n}) = \emptyset$. Then $\forall v \in \text{dom}(\mathbf{m}) : \text{frag}(v, \mathbf{m} + \mathbf{n}) = \text{frag}(v, \mathbf{m})$. In words, if $\mathbf{n}$ does not mention the VPs of $\mathbf{m}$ then it cannot influence the fragment represented by $\mathbf{m}$. This is shown by induction on the length of the longest $\text{dep}_{\mathbf{m}}$ -path (or equivalently the smallest natural number i such that $\text{dep}_{\mathbf{m}}^i = \emptyset$), which exists by the assumed acyclicity of $\mathbf{m} + \mathbf{n}$ and hence also of $\mathbf{m}$.

With that, the main claim follows easily from the definition of $\sqsubseteq$, because $\text{VP}(\mathbf{n}_1) \cap \text{VP}(\mathbf{n}_2) = \emptyset$ implies $\text{VP}(\mathbf{m}_i) \cap \text{VP}(\mathbf{n}_j) = \emptyset$ ($i \neq j$). $\square$

3 Additional SDA Operators

3.1 Subtraction

In this section we present a way of defining an “inverse” to addition. The usefulness of this operation might not be straightforward. However, as we will show, it lays the foundation for more sophisticated operations such as overriding, which is a common technique in feature-oriented software development.

For modules $\mathbf{m}$ and $\mathbf{n}$ we define the *subtraction* $\mathbf{m} - \mathbf{n}$ via restriction $|$ as

$$\mathbf{m} - \mathbf{n} =_{df} \mathbf{m} |_{\text{dom}(\mathbf{m}) - \text{dom}(\mathbf{n})} .^2$$

This spells out to

$$(\mathbf{m} - \mathbf{n})(v) =_{df} \begin{cases} \mathbf{m}(v) & \text{if } v \in (\text{dom}(\mathbf{m}) - \text{dom}(\mathbf{n})) \\ \text{undefined} & \text{otherwise} . \end{cases}$$

Among others, subtraction satisfies, for arbitrary modules $\mathbf{m}, \mathbf{n}$ and $\mathbf{p}$, the following laws. (Proofs are straightforward.)

$$\begin{array}{l|l} \begin{array}{l} \text{dom}(\mathbf{m} - \mathbf{n}) = \text{dom}(\mathbf{m}) - \text{dom}(\mathbf{n}) \\ (\mathbf{m} + \mathbf{n}) - \mathbf{p} = (\mathbf{m} - \mathbf{p}) + (\mathbf{n} - \mathbf{p}) \\ \mathbf{m} - (\mathbf{n} + \mathbf{p}) = (\mathbf{m} - \mathbf{n}) - \mathbf{p} \\ \mathbf{m} \leq \mathbf{n} \Rightarrow \mathbf{m} - \mathbf{n} = \mathbf{0} \end{array} & \begin{array}{l} \mathbf{m} - \mathbf{n} \leq \mathbf{m} \\ \mathbf{m} - \mathbf{0} = \mathbf{m} \\ \mathbf{m} - \mathbf{m} = \mathbf{0} \\ \mathbf{0} - \mathbf{n} = \mathbf{0} . \end{array} \end{array}$$

Note that the left law in the last line is only an implication, while the corresponding one for set theoretic difference is an equivalence. For the reverse direction we only have $\mathbf{m} - \mathbf{n} = \mathbf{0} \Rightarrow \text{dom}(\mathbf{m}) \subseteq \text{dom}(\mathbf{n})$. Moreover, $\mathbf{m} - \mathbf{n} \sqsubseteq \mathbf{m}$ need not hold. Subtraction is isotone in its right argument and antitone in its left, i.e.,

$$\mathbf{m} \leq \mathbf{n} \Rightarrow \mathbf{m} - \mathbf{o} \leq \mathbf{n} - \mathbf{o} , \quad \text{and} \quad \mathbf{n} \leq \mathbf{o} \Rightarrow \mathbf{m} - \mathbf{o} \leq \mathbf{m} - \mathbf{n} .$$

² Instead, one could define another subtraction operator $\ominus$ of type $(V \rightsquigarrow F(V)) \times \mathcal{P}(V) \rightarrow (V \rightsquigarrow F(V))$ by $\mathbf{m} \ominus \mathbf{U} =_{df} \mathbf{m} |_{\text{dom}(\mathbf{m}) - \mathbf{U}}$ and then set $\mathbf{m} - \mathbf{n} =_{df} \mathbf{m} \ominus \text{dom}(\mathbf{n})$.

3.2 Overriding

Ideally, modules that are composed have disjoint domains. And by using subtraction or deletion, modules can be customized. Still, object-oriented programmers are used to the notion of *overriding* or *replacing* definitions, an operation that can be defined in terms of subtraction and addition. The module $m \rightarrow n$ which results from overriding n by m is defined as

$$m \rightarrow n =_{df} m + (n - m) .$$

This replaces all assignments in n for which m also provides a value. It may destroy acyclicity. $\rightarrow$ is associative and idempotent with neutral element 0 , but not commutative.

Example 3.1. A classic example of feature interaction and product lines is the fire-and-flood control [17]. Assume a building is equipped with two systems: a fire control and a flood control. The flood control turns off the water supply as soon as sensors in the building indicate a water level is too high. It is specified by the following module.

$$\text{Flood} : \begin{cases} \text{pv} \mapsto \text{void flood()}\{ \text{pt} \} \\ \text{pt} \mapsto \text{if(isWaterHigh()) waterOff(); else waterOn();} \end{cases}$$

Fire control works in an opposite manner: as soon as sensors detect a too high temperature, sprinklers are turned on.

$$\text{Fire} : \begin{cases} \text{pw} \mapsto \text{void fire()}\{ \text{i1} \} \\ \text{i1} \mapsto \text{if(isTemperatureHigh()) spriOn(); else spriOff();} \end{cases}$$

Note that the sprinklers' functionality depends on a working water supply. Their composition/sum is $\text{FplusF} = \text{Fire} + \text{Flood}$:

$$\text{FplusF} : \begin{cases} \text{pv} \mapsto \text{void flood()}\{ \text{pt} \} \\ \text{pt} \mapsto \text{if(isWaterHigh()) waterOff(); else waterOn();} \\ \text{pw} \mapsto \text{void fire()}\{ \text{i1} \} \\ \text{i1} \mapsto \text{if(isTemperatureHigh()) spriOn(); else spriOff();} \end{cases}$$

After a fire has been detected, it must be guaranteed that the water supply is not turned off; otherwise the building would burn down. This is not guaranteed with $\text{Fire} + \text{Flood}$; it depends on a race condition determined by the order of `fire()` and `flood()` execution.

The solution is to impose one more module, representing a feature interaction FI, denoted $\text{Fire}\#\text{Flood}$ in [5], whose alterations make `Flood` and `Fire` work correctly together. This is achieved by a shared VP `pv` and the use of overriding. FI replaces `Flood`'s fragment at `pv` with FI's fragment:

FI : $\left\{ \begin{array}{l} \text{pv} \mapsto \text{void flood}()\{ \text{i2} \} \\ \text{i2} \mapsto \text{if}(\text{isTemperatureHigh}()) \text{waterOn}(); \text{ else } \{ \text{pt} \} \end{array} \right.$

The occurrence of `pt` in the fragment of `i2` may be viewed as an original call to the respective method (e.g., [8]).

The composition of all three modules, $\text{FandF} = \text{FI} \rightarrow (\text{Fire} + \text{Flood})$, yields a correct implementation:

FandF : $\left\{ \begin{array}{l} \text{pw} \mapsto \text{void fire}()\{ \text{i1} \} \\ \text{i1} \mapsto \text{if}(\text{isTemperatureHigh}()) \text{spriOn}(); \text{ else } \text{spriOff}(); \\ \text{pv} \mapsto \text{void flood}()\{ \text{i2} \} \\ \text{i2} \mapsto \text{if}(\text{isTemperatureHigh}()) \text{waterOn}(); \text{ else } \{ \text{pt} \} \\ \text{pt} \mapsto \text{if}(\text{isWaterHigh}()) \text{waterOff}(); \text{ else } \text{waterOn}(); \end{array} \right.$

□

To state further laws, we call modules m and n *compatible*, in signs $m \downarrow n$, if their fragments coincide on their shared domains. Formally,

$$m \downarrow n \Leftrightarrow_{df} \forall v \in \text{dom}(m) \cap \text{dom}(n) : m(v) = n(v) .$$

All submodules of a module are pairwise compatible with each other. It follows that the properties below are equivalent:

$$m \downarrow n \Leftrightarrow m \rightarrow n = m + n \Leftrightarrow m \rightarrow n = n \rightarrow m$$

Module addition and overriding interact by the following laws, where the left implication $\Leftarrow$ means “provided”.

$$\begin{array}{ll} m \rightarrow (n + p) = (m \rightarrow n) + (m \rightarrow p) & \text{(left distributivity)} \\ (m + n) \rightarrow p = m \rightarrow (n \rightarrow p) & \Leftarrow m \downarrow n \quad \text{(sequentialisation)} \\ (m + n) \rightarrow p = n \rightarrow p & \Leftarrow m \downarrow n \wedge m \rightarrow n = n \quad \text{(absorption I)} \\ (m + n) \rightarrow p = m \rightarrow p & \Leftarrow m \downarrow n \wedge n \rightarrow p = p \quad \text{(absorption II)} \\ m \rightarrow (n + p) = n + (m \rightarrow p) & \Leftarrow \text{dom}(m) \cap \text{dom}(n) = \emptyset \quad \text{(localisation)} \end{array}$$

The sequentialisation law means that a complex overriding may also be done by two successive simpler overridings. Absorption II, which is an immediate consequence of sequentialisation, allows simplifying a complex overriding by omitting the part that is “already there”. Finally, localisation allows propagating an overriding to the submodule it really affects.

The previous laws for overriding dealt with immediate links from VPs to fragments. We now deal with preservation of transitive links under overriding. Let dep_m^* be the reflexive, transitive closure of dep_m . For a module m and a VP $v \in \text{dom}(m)$ we define $\text{deps}(v, m) =_{df} \{w \mid v \text{ dep}_m^* w\}$, i.e., the set of VPs on which v depends transitively, plus v itself. We now want to override or extend a module n with a module m . If n does not alter the assignments of the VPs on which v transitively depends in m then the overriding/extension does not change the transitive dependence of v (cf. [23,13]):

$$\left. \begin{array}{l} \text{deps}(v, m + n) = \text{deps}(v, m) \\ \text{deps}(v, m \rightarrow n) = \text{deps}(v, n) \end{array} \right\} \Leftarrow \text{deps}(v, m) \cap \text{dom}(n) = \emptyset .$$

3.3 Solving Module Equations

As discussed in [5], it is useful to be able to solve module equations. Subtraction and its relatives enable us to do so. Suppose that m is a submodule of n , i.e., $m \leq n$ (see Equation (1)). Then the equation $m + x = n$ has $x = n - m$ as a solution. Moreover, this is the unique solution that is domain-disjoint from m .³

Example 3.2. Consider a composition $\text{comp} = a + b + c + d$ with pairwise domain-disjoint modules a, b, c, d . Then the equation $a + x + b + d = \text{comp}$ has the unique solution $x = c$ domain-disjoint from $a + b + d$. $\square$

Note that the condition $\text{dom}(m) \subseteq \text{dom}(n)$ is necessary for $m + x = n$ to be solvable, because we need to have $\text{dom}(m + x) = \text{dom}(m) \cup \text{dom}(x) = \text{dom}(n)$, which implies $\text{dom}(m) \subseteq \text{dom}(n)$. The above assumption $m \leq n$ implies that necessary condition. In fact, solvability of $m + x = n$ conversely implies $m \leq n$, since

$$m + n = m + (m + x) = (m + m) + x = m + x = n .$$

In short: $m + x = n$ is solvable iff $m \leq n$.

Next, we have a brief look at equations involving overriding. Since

$$\text{dom}(m \rightarrow n) = \text{dom}(n \rightarrow m) = \text{dom}(m) \cup \text{dom}(n) ,$$

again $\text{dom}(m) \subseteq \text{dom}(n)$ is necessary for $m \rightarrow x = n$ and $x \rightarrow m = n$ to be solvable. But by the definition of $\rightarrow$, solvability of $m \rightarrow x = n$ even implies the stronger necessary condition $m \leq n$. In this case again $x = n - m$ is the unique solution domain-disjoint from m . A closer inspection shows that the same is the case for the equation $x \rightarrow m = n$. This means that it is sufficient to restrict interest to the solution of equations involving $+$.

3.4 Transformations

In this section we sketch another extension of SDA, intended to cope with some standard techniques in software refactoring (e.g., [4,20]). Examples of such techniques are consistent renaming of methods or classes in a large software system. To stay at the same level of abstraction as before, we realize this by a mechanism for generally modifying the fragments in SDA modules.

By a *transformation* or *modification* or *refactoring* we mean a total function $T : F(V) \rightarrow F(V)$. By $T \cdot m$ we denote the *application* of T to a module m . It yields a new module defined by

$$(T \cdot m)(v) =_{df} \begin{cases} T(m(v)) & \text{if } v \in \text{dom}(m) \\ \text{undefined} & \text{otherwise} . \end{cases}$$

We use the convention that $\cdot$ binds stronger than all other operators. The following properties hold:

- (1) $\text{dom}(T \cdot m) = \text{dom}(m)$,
- (2) $T \cdot (m + n) = T \cdot m + T \cdot n \iff m \downarrow n$,
- (3) $T \cdot (m - n) = T \cdot m - T \cdot n$,
- (4) $T \cdot (m \rightarrow n) = T \cdot m \rightarrow T \cdot n$.

³ Another solution is $x = n$, since $m \leq n$ means $m + n = n$. Such solutions are uninteresting.

The proofs are straightforward calculations.

The analogue of Equation (2) is also used in the feature algebra of [1]; it entails that a transformation T is propagated and applied to all components of a composed module.

This applies, in particular, to transformations like method renaming; there T would be implemented using a global table with all the old-name/new-name associations which would be consistently applied in all submodules of the overall module under consideration. Note that, although T is supposed to be a total function on all fragments, it might well leave many of those unchanged, i.e., act as the identity on them.

Future work in this area will deal with further operators that reflect extended transformations concerning several modules, like moving methods from one module to another.

4 Using the Algebra

4.1 Projecting Out

Next, we deal with some aspects of modularity. In *classical modularity* there are physical files that define the feature modules and tools that compose them to produce a desired program. This is also known as *alternative-based variation*. Contrarily, *virtual modularity*, also known as *SYSGEN* (e.g. [14]) or *projectional variation* is a preprocessor technology. A prominent representative of this technique is the *coloring* approach of [5,15,18]. The idea is simple: every document is painted in different colors, one color per feature. A color C that appears “inside” another color D indicates a feature interaction— C changes the source of D (denoted $C\#D$ in [5]). VPs are implicit in coloring. At every point in a document where coloring changes, an implicit VP is created. Figure 1 is not only an example for variation points, fragments and modules, but also for coloring: the code is colored white, light gray and dark gray. If a feature is not needed in a product, all code colored in the corresponding color (e.g., dark gray) is *projected out*, i.e., does not show up in the final product. Since one colors the entire code base of a product line, it is possible to compute the contents of SDA modules and their VPs.

We now show how this operation can be expressed in our algebra. Assume a module m and a subset $U \subseteq \text{dom}(m)$. Projection to U is supposed to hide everything that does not correspond to a VP in U . This needs to be done in such a way that later the hidden parts can be uncovered again. Therefore, a corresponding operation should not remove the VPs outside U from m . We preserve the hidden part $n =_{df} m|_{\text{dom}(m)-U}$ and temporarily work with the module $p =_{df} [(\text{dom}(m) - U) \mapsto 0] \rightarrow m$ which masks all VPs outside U with the default fragment. To restore the original module, i.e., to switch the masked parts back on again we use that $n \rightarrow p = m$.

The difference between restriction and projecting out is that in the former case all VPs outside $\mathcal{U}$ are removed and hence become external to the resulting module so that they would be considered as “fresh” there, whereas after projection they are still “known” and hence “protected” against unintentional change, so that they can be refilled later.

4.2 Introducing Wrappers

We can use overriding for adding a wrapper to a module. Let $\mathbf{m}$ be a module and $\mathbf{v} \in \text{dom}(\mathbf{m})$ be the VP where we want to insert the wrapper. Choose a “fresh” VP $\mathbf{w} \notin \text{dom}(\mathbf{m})$ and a non-trivial wrapper fragment $\mathbf{f}$ such that $\text{VP}(\mathbf{f}) = \{\mathbf{w}\}$ and $\mathbf{w}$ occurs only once in $\mathbf{f}$. Then $\mathbf{w}$ marks the place in $\mathbf{f}$ where the original “contents” of $\mathbf{v}$ is to be put and thereby wrapped. The old contents $\mathbf{m}(\mathbf{v})$ of $\mathbf{v}$ is remembered by a link from $\mathbf{w}$ to it, so that it can be overridden by a link from $\mathbf{v}$ to the wrapper $\mathbf{f}$. Algebraically, the module with the wrapper is expressed by

$$\mathbf{m}' =_{df} ([\mathbf{v} \mapsto \mathbf{f}] \rightarrow \mathbf{m}) + [\mathbf{w} \mapsto \mathbf{m}(\mathbf{v})] .$$

An illustration is given in Figure 3.

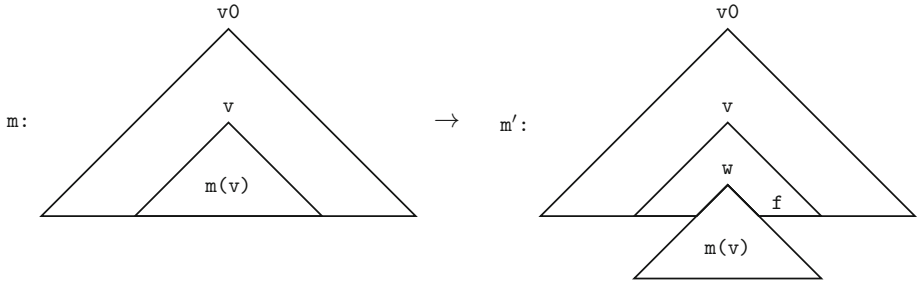


Fig. 3. Introducing a Wrapper (Splicing)

5 Small Case Study: Constructing Product Lines

There are often multiple ways how a finished product (i.e., a single fragment without further VPs) can result within a product line. This is, for example, the case when modules can be grouped into sets of features that are orthogonal to each other. Conversely, one may wish to refactor an existing monolithic system by decomposing it into directly re-usable or at least easily adaptable parts.

We illustrate this by taking up an example from [7]. Notationally we deviate a bit from that paper by not distinguishing “defining” and “applied” occurrences of VPs. All VP names start with $\textcircled{}$ (for text processing reasons); in a module every VP is assigned the text indented to the position after the VP name. If no text appears to the right of a VP this means that the default is assigned to it.

We look at software for a basic calculator. It deals with arithmetic expressions composed by the operations addition and multiplication; expressions can be

displayed and evaluated. First we present a program (or single module) **CALC**, given in [7] (see Figure 4). Formally, this is a module that has only one single root VP @0 to which the complete program is assigned, without any VPs in its text. However, to prepare the restructuring to come, we already indicate a number of further VPs, to be read as comments at this level. For operator signs **op**, such as **+** or *****, the abbreviation **?op?** stands for the string concatenation **+** "op" **+**.

CALC	
<pre>@0 abstract class Exp { @1 String print(); @2 int eval(); }</pre>	<pre>class Plus extends Exp { Exp l,r; Plus(Exp L, Exp R) {l=L; r=R;} @5 String print() { return l.print() ?+? r.print(); } @6 int eval() { return l.eval() + r.eval(); } }</pre>
<pre>class Int extends Exp { int v; Int(int a) { v=a; } @3 String print() { return v; } @4 int eval() { return v; } }</pre>	<pre>class Times extends Exp { Exp l,r; Times(Exp L, Exp R) {l=L; r=R;} @7 String print() { return l.print() ?*? r.print(); } @8 int eval() { return l.eval() * r.eval(); } }</pre>

Fig. 4. The module **CALC**

CALC contains sub-packets, which may have interest of their own. For example, a user may only want to deal with arithmetic expressions for addition. This can be achieved by grouping the parts of **CALC** into **Base**, **Plus** and **Times** as given in Figure 5 and then forming adequate sums of some of these. For instance, the product line **TPL** combines all three submodules: $\text{TPL} =_{df} \text{Base} + \text{Plus} + \text{Times}$. Using the restructuring preorder from Sect. 2.2, we have the relation $\text{CALC} \sqsubseteq \text{TPL}$ between the original code base **CALC** and **TPL**.

A second restructuring reflects that some users may only wish to evaluate expressions but not to print them. This can be accommodated by decomposing the original code base **CALC** into a **Core** module and two optional modules **Print** and **Eval**, whose definitions are given in Figure 6. A new product line $\text{CPE} =_{df} \text{Core} + \text{Print} + \text{Eval}$ combines these three submodules.

The original product line **CALC** relates to the new one by $\text{CALC} \sqsubseteq \text{CPE}$.

Finally we can form a common restructuring of **CPE** and **TPL**. This **SPL** is presented in Figure 7 and corresponds to the **EPL** matrix of [7]. We use two sets of feature names, $\mathbf{F} =_{df} \{\text{core}, \text{print}, \text{eval}\}$ and $\mathbf{G} =_{df} \{\text{base}, \text{plus}, \text{times}\}$. We use lower-case names here, since these are not module names but will serve as indices for a matrix **EPLmat**, where every entry EPLmat_{ij} ($i \in \mathbf{F}, j \in \mathbf{G}$) is a sub-module. Then setting $\text{EPL} =_{df} \sum_{i \in \mathbf{F}} \sum_{j \in \mathbf{G}} \text{EPLmat}_{ij}$ we obtain our finest restructuring of the original program **CALC**. As a consequence we have

$$\text{CPE} \sqsubseteq \text{EPL} \wedge \text{TPL} \sqsubseteq \text{EPL} .$$

TPL:		
Base:	Plus:	Times:
<pre> @0 abstract class Exp { String print(); int eval(); } class Int extends Exp { int v; Int(int a) { v=a; } String print() { return v; } int eval() { return v; } } @+ @* </pre>	<pre> @+ class Plus extends Exp { Exp l,r; Plus(Exp L, Exp R) {l=L; r=R;} String print() { return l.print() ?? r.print(); } int eval() { return l.eval() + r.eval(); } } </pre>	<pre> @* class Times extends Exp { Exp l,r; Times(Exp L, Exp R) {l=L; r=R;} String print() { return l.print() ?? r.print();} int eval() { return l.eval() * r.eval(); } } </pre>

Fig. 5. The module TPL

CPE:		
Core:	Print:	Eval:
<pre> @0 abstract class Exp { @1 @2 } class Int extends Exp { int v; Int(int a) { v=a; } @3 @4 } class Plus extends Exp { Exp l,r; Plus(Exp L, Exp R) {l=L; r=R;} @5 @6 } class Times extends Exp { Exp l,r; Times(Exp L, Exp R) {l=L; r=R;} @7 @8 } </pre>	<pre> @1 String print(); @3 String print() { return v; } @5 String print() { return l.print() ?? r.print(); } @7 String print() { return l.print() ?? r.print(); } </pre>	<pre> @2 int eval(); @4 int eval() { return v; } @6 int eval() { return l.eval() + r.eval(); } @8 int eval() { return l.eval() * r.eval(); } </pre>

Fig. 6. The module CPE

If we sum the columns of **EPLmat** we obtain the constituent modules of **CPE**, while the row sums give the constituent modules of **TPL**. Hence the small modules in **EPLmat** can be considered as the elementary features in this product line. Of course, by transitivity of $\sqsubseteq$, the original code base is an element of this product line too, i.e., $\text{CALC} \sqsubseteq \text{EPL}$; but **EPL** offers many more products.

EPLmat			
	core	print	eval
base	<pre>@0 abstract class Exp { @1 @2 } class Int extends Exp { int v; Int(int a) { v=a; } @3 @4 } @* @+</pre>	<pre>@1 String print(); eval print core @3 String print() { return v; }</pre>	<pre>@2 int eval(); @4 int eval() { return v; }</pre>
plus	<pre>@+ class Plus extends Exp { Exp l,r; Plus(Exp L, Exp R) {l=L; r=R;} @5 @6 }</pre>	<pre>@5 String print() { return l.print() ??? r.print(); }</pre>	<pre>@6 int eval() { return l.eval() + r.eval(); }</pre>
times	<pre>@* class Times extends Exp { Exp l,r; Times(Exp L, Exp R) {l=L; r=R;} @7 @8 }</pre>	<pre>@7 String print() { return l.print() ??? r.print(); }</pre>	<pre>@8 int eval() { return l.eval() * r.eval(); }</pre>

Fig. 7. The matrix EPLmat

6 Related Work

Ideas similar to those of SDA can be found in [9], where elements of UML models could be tagged with feature predicates. Given a set of selected features, an element is removed from a model if its predicate is false.

Derivatives [21] were the first identified building blocks of feature modules. Unfortunately, the mathematics of derivatives was incomplete, as composition of derivatives was not associative. This made it impossible to algebraically calculate the results of feature splitting (replacing T with $R \times S$ if T is split into features R and S) and feature merging (replacing $R \times S$ with T). CIDE [18] showed a simple and elegant way to visualize features and their interactions, resulting in the coloring algebra, which does support splitting and merging.

Other algebras for feature-based composition, such as [1,22], focus on the internal structure of modules. The algebra presented in [1] is (to our knowledge) the first that dealt with feature replication. It uses the algebraic law of *distant idempotence* (a form of idempotence where adjacency of identical features is not required). Feature composition is not commutative and there is no operation of subtraction on feature modules (called feature structure trees there). An algebra

that also captures replicas is presented in [16]. However, this algebra only works at a semantic level and cannot cope with real source code, as SDA does.

The *Choice Calculus (CC)* [25] offers an interesting and alternative approach to our work. Among the goals of CC are to integrate classical and virtual modularity, but to do so in the context of a formal programming language. Large-scale fragments can be placed in modules of their own, while small-scale fragments (suitable for annotations) can be embedded into other modules. As in coloring and `ifdef` preprocessing, variation points are implicit. The key difference between our work and CC is that the issues of classical and virtual modularity are not limited to a fixed set of programming languages.

Delta Oriented Programming (DOP) [24] is another interesting language-based approach within our field. Delta modules are qualified to be composed into a product when the corresponding `where` clause is satisfied. Such a clause is a propositional formula over features, namely the conjunction of the feature formulas that arise in coloring. Disjunction allows a single module to be reused in different contexts (rather than requiring a module to be replicated for each context). Negation seems to offer a more general way for defining alternatives. Understanding this connection will be the subject of future work. Delta modules also have `after` clauses, which specify a partial order in which to compose them. We express the composition order of modules explicitly with our overriding operator $\rightarrow\rhd$.

A project close in spirit and ideas is Kästner's *CIDE* [18,19,2]. It started out with a tool for software product line development (esp. analyzing and decomposing legacy code), following the paradigm of virtual separation of concerns. For this it offers the possibility of distinguishing fragments within the original code by different background colors. CIDE is also a compiler (and an IDE), closely related to conditional compilation with preprocessors. Additionally it analyzes the deep structure of the code and hence guarantees syntactic correctness as well as type correctness of all generated products. Hence CIDE already provides, at the tool level, some of the functionality that our algebra treats at the abstract level.

7 Conclusions and Outlook

We have presented a structured document algebra not only in concept, but also at work on some essential phenomena of large-scale software construction. We hope to have convinced the reader that SDA is both concise and precise, and comes with a number of useful laws. These can be used to construct and reason about structured modules in an algebraic fashion.

The main aim is to provide a formal basis for governing variability in large interconnected collections of documents, in particular ones that define SPLs. SDA can be used to implement tools which relieve developers from managing variability manually. In particular, it provides a basis for precise reasoning about complex and error-prone operations such as subtraction, overriding, and various refactorings. Currently, work on a pilot implementation is under way.

The algebra has interesting connections to other structures, such as the pointer algebras presented in [13,23] and to the concept of the demonic join or compatible union [3,11].

While so far SDA is presented along a concrete mathematical model, an abstraction to more general concepts like monoids, semirings and modules in the linear algebra sense is under way.

Finally, so far SDA has a very syntactic flavor, since the nature of fragments is left open. However, it is possible to work with fragments that have a semantic character, such as functions from valuations of VPs in some semantic model to a semantic value. We have sketched this in an Appendix; the elaboration of these ideas will be the subject of future investigations.

Acknowledgments. We are grateful to the anonymous referees for helpful comments and suggestions. We gratefully acknowledge support for this work by NSF grants CCF 0724979 and OCI-1148125, as well as by DFG grant MO 690/7-2. NICTA is funded by the Australian Government through the Department of Communications and the Australian Research Council through the ICT Centre of Excellence Program.

References

1. Apel, S., Lengauer, C., Möller, B., Kästner, C.: An algebraic foundation for automatic feature-based program synthesis. SCP 75(11), 1022–1047 (2010)
2. Apel, S., Batory, D., Kästner, C., Saake, G.: Feature-Oriented Software Product Lines. Springer (2013)
3. Backhouse, R., van der Woude, J.: Demonic operators and monotype factors. Mathematical Structures in Computer Science 3(4), 417–433 (1993)
4. Batory, D.: Program refactoring, program synthesis, and model-driven development. In: Adsul, B., Odersky, M. (eds.) CC 2007. LNCS, vol. 4420, pp. 156–171. Springer, Heidelberg (2007)
5. Batory, D., Höfner, P., Kim, J.: Feature interactions, products, and composition. In: Generative Programming and Component Engineering (GPCE 2011), pp. 13–22. ACM (2011)
6. Batory, D., Höfner, P., Möller, B., Zelend, A.: Features, modularity, and variation points. In: Feature-Oriented Software Development, pp. 9–16. ACM (2013)
7. Batory, D., Shepherd, C.: Product lines of product lines, Department of Computer Science, University of Texas at Austin (submitted, 2013)
8. Bergel, A., Ducasse, S., Nierstrasz, O.: Classbox/J: Controlling the scope of change in Java. SIGPLAN Not. 40(10), 177–189 (2005)
9. Czarnecki, K., Antkiewicz, M.: Mapping features to models: A template approach based on superimposed variants. In: Glück, R., Lowry, M. (eds.) GPCE 2005. LNCS, vol. 3676, pp. 422–437. Springer, Heidelberg (2005)
10. Delaware, B., Cook, W.R., Batory, D.S.: Product lines of theorems. In: Lopes, C.V., Fisher, K. (eds.) Proceedings of the 26th Annual ACM SIGPLAN Conference on Object-Oriented Programming, Systems, Languages, and Applications, OOPSLA 2011, part of SPLASH 2011, Portland, OR, USA, October 22–27, vol. 10, pp. 595–608. ACM (2011)

11. Desharnais, J., Belkhit, N., Sghaier, S.B.M., Tchier, F., Jaoua, A., Mili, A., Zaguia, N.: Embedding a demonic semilattice in a relational algebra. *Theor. Comput. Sci.* 149(2), 333–360 (1995)
12. Ducasse, S., Nierstrasz, O., Schärli, N., Wuyts, R., Black, P.: Traits: A mechanism for fine-grained reuse. *ACM Trans. Program. Lang. Syst.* 28(2), 331–388 (2006)
13. Ehm, T.: The Kleene Algebra of Nested Pointer Structures: Theory and Applications. Ph.D. thesis, Fakultät für Angewandte Informatik, Universität Augsburg (2003)
14. Gomaa, H.: Designing Software Product Lines with UML: From Use Cases to Pattern-Based Software Architectures. Addison-Wesley (2004)
15. Heidenreich, F.: Towards systematic ensuring well-formedness of software product lines. In: *Feature-Oriented Software Development*, ACM (2009)
16. Höfner, P., Khedri, R., Möller, B.: Feature algebra. In: Misra, J., Nipkow, T., Sekerinski, E. (eds.) *FM 2006*. LNCS, vol. 4085, pp. 300–315. Springer, Heidelberg (2006)
17. Kang, K., Cohen, S., Hess, J., Novak, W., Peterson, A.: Feature-oriented domain analysis (foda) feasibility study. Tech. Rep. CMU/SEI-90-TR-021, Carnegie Mellon Software Engineering Institute, Carnegie Mellon University (1990)
18. Kästner, C., Apel, S., Kuhlemann, M.: Granularity in software product lines. In: Schäfer, W., Dwyer, M., Gruhn, V. (eds.) *Conference on Software Engineering (ICSE 2008)*. ACM (2008)
19. Kästner, C.: Virtual separation of concerns: toward preprocessors 2.0. Ph.D. thesis, Otto von Guericke University Magdeburg (2010)
20. Kuhlemann, M., Batory, D., Apel, S.: Refactoring feature modules. In: Edwards, S.H., Kulczycki, G. (eds.) *ICSR 2009*. LNCS, vol. 5791, pp. 106–115. Springer, Heidelberg (2009)
21. Liu, J., Batory, D.S., Lengauer, C.: Feature oriented refactoring of legacy applications. In: Osterweil, L.J., Rombach, H.D., Soffa, M.L. (eds.) *28th International Conference on Software Engineering (ICSE 2006)*, Shanghai, China, May 20–28, pp. 112–121. ACM (2006)
22. Lopez-Herrejon, R., Batory, D., Lengauer, C.: A disciplined approach to aspect composition. In: *Partial Evaluation and Semantics-based Program Manipulation (PEPM 2006)*. ACM (2006)
23. Möller, B.: Towards pointer algebra. *SCP* 21(1), 57–90 (1993)
24. Schaefer, I., Bettini, L., Bono, V., Damiani, F., Tanzarella, N.: Delta-oriented programming of software product lines. In: Bosch, J., Lee, J. (eds.) *SPLC 2010*. LNCS, vol. 6287, pp. 77–91. Springer, Heidelberg (2010)
25. Walkingshaw, E., Erwig, M.: A calculus for modeling and implementing variation. In: *Generative Programming and Component Engineering*, pp. 132–140. ACM (2012)
26. Wirsing, M.: Algebraic specification. In: *Handbook of Theoretical Computer Science, Volume B: Formal Models and Semantics (B)*, pp. 675–788 (1990)

Appendix: Sketch of a Semantic Model of SDA

For historical reasons, in particular in honour of Martin Wirsing, we have chosen an approach based on the notions of algebraic specification [26]. Other variants are conceivable, too.

Basic Definitions. Assume a signature $\Sigma = (\mathbf{s}, \mathbf{F})$ with a sort $\mathbf{s}$ and a set $\mathbf{F}$ of operators $\mathbf{f} : \mathbf{s}^n \rightarrow \mathbf{s}$. (For simplicity we restrict ourselves to the one-sorted case; the generalisation to many sorts is straightforward.) The number n is called the *arity* of $\mathbf{f}$. An operator of arity 0 is called a *constant operator*. We also assume a set $\mathbf{V}$ of *variation points* such that $\mathbf{V}$ is disjoint from the set of constant operators of Σ (in logic the elements of $\mathbf{V}$ would be called *variables*). The set $\mathbf{T}(\Sigma, \mathbf{V})$ of *terms* over Σ and $\mathbf{V}$ is defined as usual. If $\mathbf{VP}(\mathbf{t}) = \emptyset$ then $\mathbf{t}$ is called *closed* or a *ground term*. For describing SDA modules we assume a signature with two special constant symbols 0, $\perp$.

A Σ -*algebra* is a pair $\mathbf{A} = (\mathbf{s}^{\mathbf{A}}, \mathbf{F}^{\mathbf{A}})$ where $\mathbf{s}^{\mathbf{A}}$ is a nonempty *carrier set* and $\mathbf{F}^{\mathbf{A}} = \{\mathbf{f}^{\mathbf{A}} \mid \mathbf{f} \in \mathbf{F}\}$ is a set of *operations* $\mathbf{f}^{\mathbf{A}} : (\mathbf{s}^{\mathbf{A}})^n \rightarrow \mathbf{s}^{\mathbf{A}}$ associated with the $\mathbf{f} : \mathbf{s}^n \rightarrow \mathbf{s} \in \mathbf{F}$. The set $\mathbf{T}(\Sigma, \mathbf{V})$ can be made into a Σ -algebra in the usual way. A *valuation* of $\mathbf{V}$ in a Σ -algebra $\mathbf{A}$ (also called an *environment*) is a partial function $\mathbf{e} : \mathbf{V} \rightsquigarrow \mathbf{s}^{\mathbf{A}}$. The set of all environments is denoted by $\mathcal{E}^{\mathbf{A}}$.

Syntactic Modules. A *syntactic fragment* simply is an element of $\mathbf{T}(\Sigma, \mathbf{V})$. By this, all VPs are fragments, too. A *syntactic module* is an environment $\mathbf{m} : \mathbf{V} \rightsquigarrow \mathbf{T}(\Sigma, \mathbf{V})$ from $\mathcal{E}^{\mathbf{T}(\Sigma, \mathbf{V})}$. The function $\mathbf{single_fill}(\mathbf{t}, \mathbf{m})$ is defined inductively as syntactic substitution:

1. If $\mathbf{t} = \mathbf{v} \in \mathbf{V}$ is a VP then

$$\mathbf{single_fill}(\mathbf{t}, \mathbf{m}) =_{df} \begin{cases} \mathbf{e}(\mathbf{v}) & \text{if } \mathbf{v} \in \mathbf{dom}(\mathbf{e}) , \\ \mathbf{v} & \text{otherwise .} \end{cases}$$

2. If $\mathbf{t} = \mathbf{f}(\mathbf{t}_1, \dots, \mathbf{t}_n)$ then

$$\mathbf{single_fill}(\mathbf{t}, \mathbf{m}) =_{df} \mathbf{f}(\mathbf{single_fill}(\mathbf{t}_1, \mathbf{m}), \dots, \mathbf{single_fill}(\mathbf{t}_n, \mathbf{m})) .$$

For a constant operator $\mathbf{f}$ hence $\mathbf{single_fill}(\mathbf{t}, \mathbf{m}) = \mathbf{f}$.

Term Evaluation. For a Σ -algebra $\mathbf{A}$ we define inductively the *evaluation* $\llbracket _ \rrbracket : \mathbf{T}(\Sigma, \mathbf{V}) \rightarrow (\mathcal{E}^{\mathbf{A}} \rightsquigarrow \mathbf{s}^{\mathbf{A}})$ that assigns to every term a value using an environment, if possible.

- For VP $\mathbf{v}$ we set

$$\llbracket \mathbf{v} \rrbracket(\mathbf{e}) =_{df} \begin{cases} \mathbf{e}(\mathbf{v}) & \text{if } \mathbf{v} \in \mathbf{dom}(\mathbf{e}) , \\ \text{undefined} & \text{otherwise .} \end{cases}$$

- For other terms we set

$$\llbracket \mathbf{f}(\mathbf{t}_1, \dots, \mathbf{t}_n) \rrbracket(\mathbf{e}) =_{df} \begin{cases} \mathbf{f}^{\mathbf{A}}(\mathbf{u}_1, \dots, \mathbf{u}_n) & \text{if all } \llbracket \mathbf{t}_i \rrbracket(\mathbf{e}) \text{ are defined and } \mathbf{u}_i = \llbracket \mathbf{t}_i \rrbracket(\mathbf{e}) , \\ \text{undefined} & \text{otherwise .} \end{cases}$$

Hence for a constant symbol $\mathbf{f}$ we have $\llbracket \mathbf{f} \rrbracket(\mathbf{e}) = \mathbf{f}^{\mathbf{A}}$ for all $\mathbf{e}$.

This definition entails what is called the *Coincidence Lemma* in logic; the proof is a straightforward induction.

Lemma 7.1. *If two environments e, e' agree on the VPs of a term t , i.e., if $e|_{VP(t)} = e'|_{VP(t)}$, then $\llbracket t \rrbracket(e) = \llbracket t \rrbracket(e')$.*

Semantic Modules. The idea is now to make the mapping $\llbracket t \rrbracket : \mathcal{E}^A \rightsquigarrow \mathbf{s}^A$ somehow into a *semantic fragment* corresponding to the syntactic object t and to define a semantic module as a partial function from VPs to semantic fragments. However, this is not entirely straightforward, since we need to have VP information in some way to still be able to talk about cycle-freeness of modules. The solution proposed here is to enrich a semantic fragment by a set of VPs it “administers”. A *semantic fragment* over a set V of VPs and a Σ -algebra A is a pair (W, g) where $W \subseteq V$ and $g : \mathcal{E}^A \rightsquigarrow \mathbf{s}^A$ satisfies the *coincidence property* on W :

$$\forall e, e' : e|_W = e'|_W \Rightarrow g(e) = g(e') .$$

We set $VP(W, g) =_{df} W$ and $ev(W, g) =_{df} g$. The dependence relation for semantic modules uses this definition of VP.

Every VP v can be made into the fragment $(\{v\}, \lambda e. 0^A)$.

A *semantic module* is a partial function from V to the set of semantic fragments. We can translate every syntactic module m into a semantic one called $\hat{m}$ by setting $\hat{m}(v) =_{df} (VP(m(v)), \llbracket m(v) \rrbracket)$. If m is cycle-free then so is $\hat{m}$.

The Single Fill Function for Semantic Modules. To make SDA work for semantic modules we have to define the `single_fill` function suitably:

$$\text{single_fill}((W, g), m) =_{df} ((W - \text{dom}(m)) \cup Z, \lambda e. g(e')) ,$$

where $Z =_{df} \bigcup_{v \in \text{dom}(m)} VP(m(v))$ and

$$e'(v) =_{df} \begin{cases} m(v)(e) & \text{if } v \in \text{dom}(m) , \\ e(v) & \text{otherwise} . \end{cases}$$