

Теории типов — лекционные заметки

Alexander Gryzlov¹ and Alexander Kuklev^{2,3} <a@kuklev.com>

¹IMDEA Software Institute, Spain

²Radboud University Nijmegen, Software Science

³JetBrains Research

1. Вводная лекция

Добро пожаловать на курс по теориям типов!

Первым делом нужно рассказать, что вообще такое теории типов. А для этого нужно сперва немного поговорить о языках и системах записи, начиная с естественных языков, таких как тот, на котором написан этот текст.

1.1. Специализированные языки

Естественные языки очень универсальны. Кроме собственно коммуникации между людьми, они позволяют описывать факты и действия, доказательства и алгоритмы, абстрактные концепции и материальные устройства, идеализированные физические системы и реальные экспериментальные установки. Вместе с тем, естественные языки имеют фундаментальный недостаток, особенно критичный во времена LLM: они допускают разумно звучащую белиберду. Словесные описания длинных алгоритмов, доказательств или спецификаций бизнес-процессов/бизнес-логики зачастую содержат незаметные или намеренно замаскированные уязвимости, ошибки, неоднозначности и упущения. Проблемы такого рода продолжают находить спустя годы и десятилетия в текстах, внимательно вычитанных сотнями экспертов.

- В канадском споре Rogers Comm. и Aliant Telecom смысл 14-страничного договора зависел от запятой в положении о сроке действия и прекращении договора. Изначально спор мог стоить Rogers около 1 млн долларов; CRTC сначала приняла одну интерпретацию, затем, после анализа французской версии соглашения, изменила решение.
- Выдающийся математик Владимир Воеводский стал лауреатом Филдсовской премии за свои работы в области теории мотивов. Частью фундамента этой работы был технический аппарат, описанный им в статье “Cohomological Theory of Presheaves with Transfers” в 1992 году. У статьи сотни цитирований; начиная с 1993 года несколько групп математиков изучали эту статью на семинарах и использовали её в своей работе, но годами никто не замечал ошибки. Её заметил сам автор — спустя целых восемь лет, и ошибка была существенной. Не только доказательство одной из ключевых лемм содержало ошибку — сама формулировка леммы не подлежала спасению. Потратив ещё шесть лет, Воеводский сумел заменить её более слабой и более сложной леммой, которой всё же хватало для всех приложений. Именно этот опыт стал причиной его интереса к машинной проверке доказательств и прорывным идеям, о которых мы будем много говорить в ходе этого курса.

Справляться с фундаментальным недостатком естественных языков позволяют специализированные формальные языки, устраняющих неоднозначности и навязывающих структурированность и систематичность изложения. Неоднозначность становится невыразимой, а лакуны и одумки становятся явными. Специализированные языки должны быть пригодными для алгоритмической обработки и в то же время максимально удобными для чтения и понимания человеком — удобнее естественного языка для квалифицированного читателя.

Немало именно таких языков появилось задолго до появления компьютеров. Один из старейших примеров — позиционная запись чисел цифрами:

- записать, прочитать и понять в голове число 2524 значительно проще, чем «две тысячи пятьсот двадцать четыре»,
- запись цифрами автоматически исключает числа-одумки («двадцать восемь, двадцать девять, двадцать десять»),
- удачная специализированная запись открывает новые возможности — арифметические операции в столбик. Это как раз пример “алгоритмической обработки”, ставшей возможной благодаря удачному специализированному способу записи.

Ещё один важный пример — химическая запись (C_2H_5OH , anyone?), значительно поспособствовавшая в своё время пониманию химических реакций. А вот запись чисел римскими цифрами — это пример-предостережение, как специализированный язык может погрязнуть в избыточной сложности в результате отсутствия какой-либо идеи. В данном случае — концепции нуля.

1.2. Языки программирования

Наверняка, самый знакомый вам пример — языки программирования. Программы на них — это просто конечные последовательности символов заданного наперёд алфавита. Но не любые, а лишь те, которые допускает компилятор/интерпретатор языка. А семантика языка программирования (то есть исчерпывающее смысловое значение текстов на нём) тоже в конечном итоге определяется интерпретатором/компилятором и поведением машины, на которой программа исполняется.

Для многих языков программирования написана книга под названием “стандарт языка”, где описана формальная грамматика языка (в терминах строгих синтаксических правил) и семантика языка, как правило, в терминах того, как программа должна в точности исполняться на идеализированной/стандартизированной машине.¹ Отметим, однако, что почти для всех реальных языков программирования стандарт языка — это больше декларация о намерениях, чем исчерпывающее описание.

Исключение составляют языки программирования, возникшие не из практических нужд, а из математических абстракций — в частности, языки на базе комбинаторной логики/ λ -исчисления. Почти все такие языки программирования низкоуровневые, то есть больше похожи на машинный код — невозможно понять смысл программы нетривиальной длины, читая код; результат работы программы (даже если это просто число) как правило нуждается в расшифровке, а для написания нетривиальных программ нужен ещё слой сокращений-макросов, своего рода вторичный язык.

¹Это называется операционной семантикой. Существуют и другие подходы, среди которых нужно отдельно отметить денотационную и аксиоматическую семантику. О том и другом мы ещё подробно поговорим в следующих лекциях.

1.3. System T: Особенный язык программирования

Первый высокоуровневый “математический” язык программирования был побочным результатом математической работы 1958 года Курта Гёделя о непротиворечивости арифметики первого порядка. Автор не планировал разрабатывать язык программирования, не интересовался машинной имплементацией и совершенно не задумывался о практическом удобстве — устройство языка диктовалось соображениями логической структуры арифметики.

Тем не менее, результатом стал высокоуровневый строго-типизированный язык программирования System T на базе λ -исчисления и самой простой формы рекурсии, где числа — это числа, их последовательности — это последовательности, а программы напрямую работают с такими объектами (без постоянного кодирования и декодирования) и используют достаточно высокоуровневые примитивы, чтобы нетривиальные алгоритмы в самом деле можно было писать руками и читать глазами.

Удивительная особенность этого языка программирования — поразительная элегантность и простота его описания. В то время как стандарт языка C — это книга на 284 страницы, System T полностью описывается небольшим набором правил, который легко вмещается на полстраницы (см. ниже) и исчерпывающе описывает синтаксис, семантику и систему типов языка.

Правила отталкиваются именно от системы типов: они описывают, какие типы данных встречаются в языке, и для каждого типа даётся исчерпывающий набор операторов порождения и разбора значений этого типа; кроме того, описывается действие операторов разбора на операторы порождения.

$$\begin{array}{c}
 \frac{}{\Gamma, x : T \vdash x : T}(\text{var}) \qquad \frac{}{\text{Unit}}(\text{unit}) \qquad \frac{}{\Gamma \vdash () : \text{Unit}} \\
 \\
 \frac{X \quad Y}{X \rightarrow Y}(\text{func}) \qquad \frac{\Gamma, x : X \vdash f : Y}{\Gamma \vdash \lambda(x : X) f : X \rightarrow Y} \qquad \frac{\Gamma \vdash f : X \rightarrow Y \quad \Gamma \vdash x : X}{\Gamma \vdash f x : Y} \\
 \\
 \frac{X \quad Y}{X \times Y}(\text{pair}) \qquad \frac{\Gamma \vdash x : X \quad \Gamma \vdash y : Y}{\langle x, y \rangle : X \times Y} \qquad \frac{\Gamma \vdash p : X \times Y}{\Gamma \vdash p_{\text{fst}} : X} \qquad \frac{\Gamma \vdash p : X \times Y}{\Gamma \vdash p_{\text{snd}} : Y} \\
 \\
 \frac{T}{T^*}(\text{list}) \qquad \frac{}{\Gamma \vdash \text{nil} : T^*} \qquad \frac{\Gamma \vdash \text{head} : T \quad \Gamma \vdash \text{tail} : T^*}{\Gamma \vdash \text{head}, \text{tail} : T^*} \\
 \\
 \frac{\Gamma \vdash \text{list} : X^* \quad \Gamma \vdash \text{init} : Y \quad \Gamma, \text{head} : X, \text{tail} : X^*, \text{prev} : Y \vdash \text{step} : Y}{\Gamma \vdash \text{reduce}(\text{list})(\text{step})(\text{init}) : Y} \\
 \\
 (\lambda(x : X) f(x)) x' \rightsquigarrow f(x') \qquad \langle x, y \rangle_{\text{fst}} \rightsquigarrow x \qquad \langle x, y \rangle_{\text{snd}} \rightsquigarrow y \\
 \\
 \text{reduce}(\text{head}, \text{tail})(\text{step})(\text{init}) \rightsquigarrow \text{step}(\text{head})(\text{tail})(\text{reduce}(\text{tail})(\text{step})(\text{init})) \\
 \\
 \text{reduce}(\text{nil})(\text{step})(\text{init}) \rightsquigarrow \text{init}
 \end{array}$$

Такой подход называется типотеоретическим описанием языка, а языки, описанные таким образом, называются теориями типов (в узком смысле). Математическая дисциплина, изучающая их, тоже называется теорией типов (в широком смысле). Типотеоретическое описание языка не только очень элегантно и экономично, но ещё и очень удобно для того, чтобы математически рассуждать о свойствах программ, доказывать или опровергать их свойства.

Большую часть языков программирования (C, Java) невозможно описать так экономно, но можно разрабатывать языки программирования, отталкиваясь

от возможностей типотеоретического описания. Разработчик всё ещё исходит из практических нужд программирования, но вынужден действовать в очень строгих рамках возможного. Эти рамки очень часто, *more often than not*, подсказывают элегантные дизайнерские решения. Иногда они не просто элегантные, а вообще единственно возможные. Как и в математике, стирается грань между исследованием и разработкой. Разработчик языка одновременно “придумывает” новые конструкции языка и “открывает” правильный способ их реализации, вытекающий из самой природы вещей.

1.4. Гораздо больше, чем языки программирования

Хотя теории типов неразрывно связаны с *computer science*, их область применимости гораздо шире, чем программирование. Собственно, успешных мейнстримных языков программирования с исчерпывающим типотеоретическим описанием до сих пор нет, хотя мы и приближаемся к этому.

Теория типов произвела грандиозную революцию в области машинопроверяемых доказательств и оснований математики. Типотеоретический подход подсказал элегантные способы улучшить языки для формулирования и доказательства математических теорем. Эти языки, называемые исчислениями построений/конструктивными теориями типов, всё ещё родственны языкам программирования, а точнее включают в себя тотальные функциональные языки программирования наряду с языком доказательств.

Однако теории типов могут быть совершенно не связанными с программированием. Так, в 2018-2023 годах были получены типотеоретические описания ω -категорий и родственных объектов, что привело к существенному прогрессу в соответствующем разделе математики.

С другой стороны, основанные на теории типов языки программирования могут быть обобщены до языков описания взаимодействующих систем (см. *KeymaeraX*,² *VeriDrone*,³ *SpaceEx*⁴). В том числе гибридных систем, где часть компонентов действует согласно детерминистическому алгоритму, а часть лишь удовлетворяет определённым (непрерывным и/или вероятностным) моделям поведения, описываемым в терминах кусочно-дифференциальных уравнений. Такие языки уже сейчас служат основой для формальных доказательств безопасности беспилотных автомобилей и летательных аппаратов, газопроводов и реакторов.

Дальнейшим обобщением являются языки, способные описывать экспериментальные установки и физические системы, в которых могут играть роль квантовые эффекты — область, где естественные языки как раз теряют применимость. В последние пять лет достигнут огромный прогресс в области экспериментальных теорий типов, естественным образом описывающих квантовые теории поля и теории струн (*Myers et al.*, 2024; *Sati and Schreiber*, 2026). Такие теории призваны предоставить язык, позволяющий достичь концептуального понимания в проблеме измерения, проблеме локального реализма и, в конечном итоге, проблеме времени.

²<https://keymaerax.org/>

³<https://ucsd-pl.github.io/veridrone/>

⁴<http://spaceex.imag.fr/>

2. Языки и алгебраические теории

Типизированные языки задаются грамматическими правилами вида:

$$\frac{s_1 : T_1 \quad \dots \quad s_n : T_n}{e : T}$$

Грамматическое правило описывает, как выражения складываются из подвыражений. Под горизонтальной чертой указывается описываемое выражение и – через двоеточие – его тип. Над горизонтальной чертой указываются несколько (их может быть ноль) подвыражений, используемых в выражении. В естественных языках типы — это части речи (существительное) или типы функциональных групп (напр. “новый алый шарф” — номинальная группа).⁵

В качестве примера давайте рассмотрим язык алгебраических выражений со скалярами и векторами из евклидова векторного пространства над ними:

$$\begin{array}{c} \frac{a : \text{Scalar} \quad v : \text{Vector}}{av : \text{Vector}} \qquad \frac{v : \text{Vector} \quad u : \text{Vector}}{v \cdot u : \text{Scalar}} \\[10pt] \frac{v : \text{Vector} \quad u : \text{Vector}}{v + u : \text{Vector}} \qquad \frac{v : \text{Vector}}{-v : \text{Vector}} \qquad \frac{}{0 : \text{Vector}} \\[10pt] \frac{a : \text{Scalar} \quad b : \text{Scalar}}{a + b : \text{Scalar}} \qquad \frac{a : \text{Scalar}}{-a : \text{Scalar}} \qquad \frac{}{0 : \text{Scalar}} \\[10pt] \frac{a : \text{Scalar} \quad b : \text{Scalar}}{ab : \text{Scalar}} \qquad \frac{}{1 : \text{Scalar}} \end{array}$$

В типизированных языках выражения отождествляются со своими деревьями вывода. Поэтому правила не содержат скобок, приоритизации и ассоциативности операторов. Преобразование выражений в текст и обратно (парсинг текста в синтаксические деревья) — это отдельная область, называемая теорией формальных языков, и её мы тут рассматривать не будем.

Типизированные языки могут иметь бесконечные ряды типов и правил. Характерный пример — нетипизированное λ -исчисление, где тип выражения определяется количеством свободных переменных в нём, и стало быть, типов бесконечно много. Грамматические правила там получаются из нижеследующих схем путём подстановки натуральных чисел в качестве индексов:

$$\frac{f : \text{Term}_n \quad x : \text{Term}_n}{fx : \text{Term}_n}(\text{app}) \quad \frac{t : \text{Term}_{n+i+1}}{\lambda_i t : \text{Term}_{n+i}}(\text{abs}) \quad \frac{}{x_n : \text{Term}_n}(\text{var}) \quad \frac{t : \text{Term}_n}{t : \text{Term}_{n+1}}(\text{sub})$$

Использование нумерованных переменных не слишком удобно на практике: людям удобнее воспринимать именованные переменные. В связи с этим был разработан особый формат записи грамматических правил, называемый естественным выводом. Вместо того чтобы писать в типе выражения количество переменных, перед значением записывается контекст — список именованных переменных, используемых в нём.

$$\begin{array}{c} \frac{\Gamma \vdash f : \text{Term} \quad \Gamma \vdash x : \text{Term}}{\Gamma \vdash fx : \text{Term}}(\text{app}) \qquad \frac{\Gamma, x \vdash t : \text{Term}}{\Gamma \vdash (x \mapsto t) : \text{Term}}(\text{abs}) \\[10pt] \frac{}{x \vdash x : \text{Term}}(\text{var}) \qquad \frac{\Gamma \vdash t : \text{Term}}{\Gamma, x \vdash t : \text{Term}}(\text{sub}) \end{array}$$

⁵Типизированные языки используются и в лингвистике. См. Категориальные грамматики Ламбека — это формальные системы для математически строгого анализа синтаксиса и семантики естественных языков.

Для полноты картины давайте приведём ещё язык теории множеств:

$$\begin{array}{c}
\frac{}{x, y \vdash x = y : \text{Prop}} (\text{eq}) \qquad \frac{}{x, y \vdash x \in y : \text{Prop}} (\text{in}) \qquad \frac{\Gamma \vdash p : \text{Prop}}{\Gamma, x \vdash p : \text{Prop}} \text{sub} \\
\\
\frac{\Gamma, x \vdash p : \text{Prop}}{\Gamma \vdash \exists(x) p : \text{Prop}} (\text{exists}) \qquad \frac{\Gamma, x \vdash p : \text{Prop}}{\Gamma \vdash \forall(x) p : \text{Prop}} (\text{forall}) \\
\\
\frac{\Gamma \vdash p : \text{Prop} \quad \Gamma \vdash q : \text{Prop}}{\Gamma \vdash p \wedge q : \text{Prop}} (\text{and}) \qquad \frac{\Gamma \vdash p : \text{Prop} \quad \Gamma \vdash q : \text{Prop}}{\Gamma \vdash p \rightarrow q : \text{Prop}} (\text{implies}) \\
\\
\frac{\Gamma \vdash p : \text{Prop} \quad \Gamma \vdash q : \text{Prop}}{\Gamma \vdash p \vee q : \text{Prop}} (\text{or}) \qquad \frac{\Gamma \vdash p : \text{Prop}}{\Gamma \vdash \neg p : \text{Prop}} (\text{not})
\end{array}$$

Отметим, что естественная дедукция — всего лишь способ записи. Имеются механические способы перевода правил естественной дедукции в обычные синтаксические правила — de Bruijn и co-de Bruijn translations.

Естественная дедукция значительно упрощает описание типизированных λ -исчислений, языков аксиоматических теорий и, совершенно особенно, языков доказательств, для которых она, собственно, и была придумана.

2.1. Алгебраические теории

Грамматика языка ещё не описывает его семантики, однако во некоторых очень важных случаях для исчерпывающего описания семантики достаточно добавить в язык тождества — правила о том, в каких случаях выражения равны между собой. Тождества тоже записываются как правила с горизонтальной чертой. Сверху — переменные, снизу — равенство двух выражений. Язык вместе с набором тождеств называется алгебраической теорией.

Например, множества с одной ассоциативной операцией (например, целые числа со сложением) суть модели следующей алгебраической теории:

$$\frac{x : S \quad y : S}{xy : S} \qquad \frac{x : S \quad y : S \quad z : S}{(xy)z = x(yz)}$$

Модели этой теории, кстати, называются в математике полугруппами. Если ещё добавить нейтральный элемент — получатся моноиды. Умножение матриц тоже ассоциативная операция, но тут нам понадобится типизированная теория, где типы характеризуют размеры матриц.

$$\frac{x : M_n^m \quad y : M_k^n}{\Gamma \vdash xy : M_k^m} \qquad (xy)z = x(yz)$$

Такие типизированные моноиды называются категориями, и являются одной из самых часто встречающихся и богатых структур в математике. Разумеется, в общем случае индексы m и n не обязаны быть положительными целыми числами, а могут быть представителями любого множества (точнее, класса).

Лупы, группы, кольца, решётки, алгебры Ли — всё это также алгебраические теории. Модели алгебраических теорий образуют категории, обладающие массой элегантных свойств. Если тот или иной класс математических объектов удаётся охарактеризовать в терминах алгебраических теорий, это даёт глубокое структурное понимание объектов и широчайший инструментарий математических средств для работы с ним. Среди прочего, для каждой алгебраической теории существует уникальная универсальная модель, а для каждого множества порождающих

существует свободно-порождённая ими модель. Так, например, свободный моноид на n порождающих — это просто множество строк алфавита из n букв вместе с операцией конкатенации.

Ценные свойства алгебраических теорий вытекают из того, что операция построения свободно-порождённой модели является монадой (что бы это ни значило), и порождённая этой монадой категория — в точности категория всех моделей алгебраической теории с гомоморфизмами между ними.

Когда мы говорили о типизированных языках, мы требовали, чтобы каждое правило вывода содержало лишь конечное число посылок (подвыражений), а сами грамматические правила были эффективно перечислимы, т.е. существовала компьютерная программа, которая последовательно генерирует все правила. Эти ограничения гарантируют, что выражения языка тоже эффективно перечислимы. Это соответствует интуитивным представлениям о том, что такое язык — перечислимость выражений означает, что в конечном итоге все они могут быть записаны как конечный текст на конечном алфавите.

Тем не менее, ценные математические свойства алгебраических теорий сохраняются и в случае, если допускаются непечислимо большие языки. Первый пример — параметризованные языки. Мы можем взять язык, где для каждого вещественного числа $s : \mathbb{R}$ имеется операция скалярного умножения

$$\frac{v : \text{Vector} \quad u : \text{Vector}}{v + u : \text{Vector}} \qquad \frac{v : \text{Vector}}{sv : \text{Vector}} \text{ for each } s : \mathbb{R}$$

Нам потребуются также параметризованные тождества, гарантирующие ассоциативность и дистрибутивность скалярного произведения. Модели этой теории — в точности векторные пространства над вещественными числами.

Другой пример — теории с оператором, принимающим бесконечное число аргументов. Скажем, оператор $\lim(a_n)$ предела последовательности.⁶ Модели такой теории — компактные Хаусдорфовы пространства.

В обоих случаях язык теории оказывается несчётным, но это не мешает построить её монаду. Такие теории называются инфинитарно-алгебраическими.

2.2. Внутренние языки категорий

В 1970ые годы было обнаружено, что важный класс категорий (декартово замкнутые категории) — суть в точности модели алгебраической теории STLC (просто типизированного λ -исчисления). Таким образом, все построения, доступные в декартово-замкнутых категориях благодаря их структуре, могут быть описаны на языке STLC, а все их свойства вытекают из тождеств STLC. Так сформировалось понятие внутреннего языка класса категорий — алгебраической теории, моделями которой являются категории этого класса.

Вскоре [Seely \(1984\)](#) сформулировал аналогичный результат для локально декартово замкнутых категорий, см. [Clairambault and Dybjer \(2014\)](#) для полного доказательства. Проблема в том, что их “внутренним языком” является экстенциональная теория типов $\lambda_{\text{PTS}^\equiv}$, не являющаяся алгебраической теорией. Тут нужно перейти от строгих категорий к ∞ -категориям, где свойства композиции соблюдаются не строго, а с точностью до эквивалентности. У локально декартово-замкнутых ∞ -категорий в самом деле есть внутренний язык — унивалентная теория типов $\lambda_{\text{PTS}^\approx}$ ([Kapulkin, 2017](#)).

⁶Для полной общности нужно использовать ультрафильтры вместо последовательностей.

В ходе этого курса мы будем говорить о внутренних языках других классов категорий, важных с типотейоретической точки зрения, см. рис. 1.

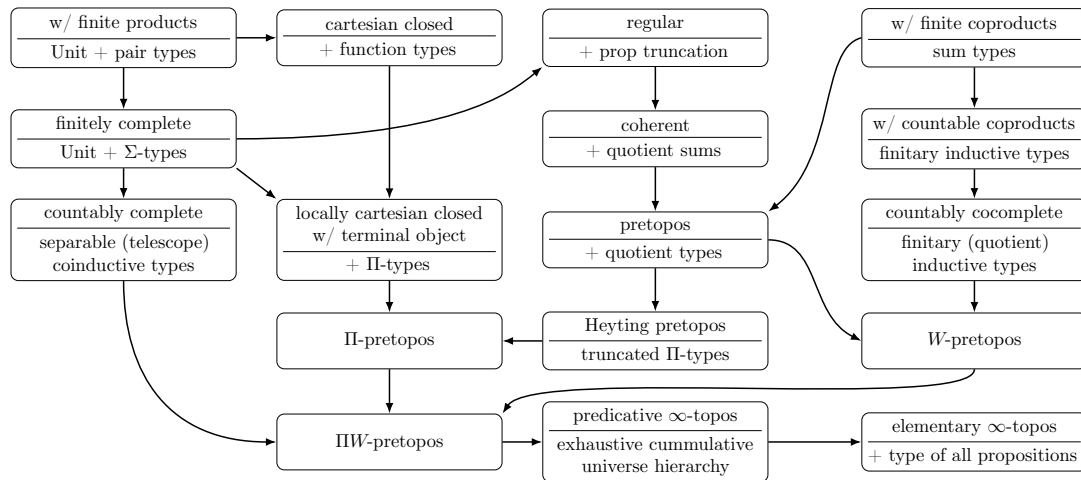


Рис. 1: Основные классы типотейоретически важных ∞ -категорий

Чтобы быть совсем точными, отметим, что предикативные и элементарные ∞ -топосы обладают внутренними языками в более слабом смысле. Все эти ∞ -топосы являются моделями соответствующей теории, но не все модели являются такими ∞ -топосами. Однако ими являются все конечно-порождённые модели, и в том числе универсальный элементарный ∞ -топос.

2.3. Универсальный элементарный ∞ -топос

С начала 1920ых годов в качестве общей базы для всей математики использовалась теория множеств Цермело-Френкеля ZFC. К сожалению, эта теория не является в каком-либо смысле канонической — у неё есть масса вариаций или аналогов, которые как будто ничем не хуже, но и не сильно лучше.

В середине 20го века начала развиваться теория категорий — дисциплина, посвящённая главным образом абстрагированию квинтэссенции конкретных свойств объектов. Ключевым свойством ZFC, делающим её пригодной на роль общей базы, является возможность конструировать в ней модели аксиоматических теорий. Это свойство было абстрагировано — в конце 1960ых был охарактеризован класс категорий, обладающий этим свойством. Они называются элементарными топосами, и их внутренний язык (то есть язык универсального элементарного ∞ -топоса) является каноническим языком математики.

Поиск этого языка — путешествие длиной в 40 лет, близкое к завершению. На данный момент язык реализован в пруф-ассистенте Narya,⁷ за исключением одной технической недоделки⁸ и одной, правда весьма ключевой, теоремы.⁹

⁷<https://narya.readthedocs.io/en/latest/>

⁸Не реализованы индуктивные семейства — их реализация там весьма сложна технически.

⁹Доказательство фибрантности унивалентных вселенных.

3. STLC и System T — подробный разбор

Во вводной лекции System T появился как первый высокоуровневый «математический» язык программирования, выросший из работы Гёделя о непротиворечивости арифметики. Теперь мы рассмотрим его как тотальный язык, построим для него нормализатор, а затем вернёмся к его исходной роли целевого языка для трансляции арифметических доказательств.

Каждая программа System T завершается (это и есть свойство тотальности). Благодаря этому программы всегда можно полностью нормализовать, а затем сравнивать и анализировать как конечные синтаксические деревья.

Весь материал лекции формализован в Rosq: интерпретация Диалектика находится в репозитории `dialectica-system-t`,¹⁰ а нормализация в его подмодуле `nbe-system-t`,¹¹ следующем второй главе диссертации [Abel \(2013\)](#). Ядро STLC можно отдельно посмотреть в реализации на Idris.¹²

3.1. Конечные типы данных, суммы, пары, функции

Как мы подчеркнули в первой лекции, экономичное описание языков программирования отталкивается от доступных в них типов данных. Самый простой содержательный тип данных — бит, булево значение. Условимся обозначать этот тип $\mathbb{B}$, а его канонические значения `tt` и `ff`. Для всякого типа данных нужно указать синтаксис порождения значений и синтаксис их разбора. Синтаксис порождения просто описывает возможные канонические значения:

$$\frac{}{\Gamma \vdash \text{tt} : \mathbb{B}} \qquad \frac{}{\Gamma \vdash \text{ff} : \mathbb{B}}$$

Синтаксис разбора булевых значений — привычный оператор `if`:

$$\frac{\Gamma \vdash b : \mathbb{B} \quad \Gamma \vdash \text{then} : T \quad \Gamma \vdash \text{else} : T}{\Gamma \vdash \text{if}(b) \text{ then else} : T}$$

То, как “работает” оператор, задают правила вычисления (= β -reduction):

$$\text{if}(\text{tt}) \text{ then else} \rightsquigarrow \text{then} \qquad \text{if}(\text{ff}) \text{ then else} \rightsquigarrow \text{else}$$

Посторонние значения исключаются правилом уникальности (= η -expansion):

$$\frac{\Gamma \vdash b : \mathbb{B}}{b \equiv \text{if}(b) \text{ tt ff}}$$

Стандартные действия над натуральными числами — $n+m$, $n \cdot m$ и n^m — мотивированы операциями над конечными множествами. Так n^m это количество отображений из множества с m элементами во множество с n элементами; $n \cdot m$ это количество (упорядоченных) пар, где первый элемент из множества с n элементами, а второй — из множества с m элементами. Наконец, сумма $n + m$, это размер размеченного объединения множеств с n и m элементами соответственно. Давайте опишем соответствующие операции на типах.

Если у нас есть два типа данных X и Y , можно задать их размеченное объединение

¹⁰<https://github.com/clayrat/dialectica-system-t/>

¹¹<https://github.com/clayrat/nbe-system-t/>

¹²<https://github.com/sequents/code/tree/master/src/Lambda>

$(X + Y)$ — обобщение типа $\mathbb{B}$:

$$\begin{array}{c}
\frac{\Gamma \vdash x : X}{\Gamma \vdash \text{inl}(x) : (X + Y)} \qquad \frac{y : Y}{\Gamma \vdash \text{inr}(y) : (X + Y)} \\
\\
\frac{\Gamma \vdash e : (X + Y) \quad \Gamma, x : X \vdash \text{left} : T \quad \Gamma, y : Y \vdash \text{right} : T}{\Gamma \vdash \text{match}(e) \{ \text{inl } x \mapsto \text{left}; \text{inr } y \mapsto \text{right} \} : T} \\
\\
\text{match}(\text{inl } x) \{ \text{inl}(x) \mapsto \text{left}; \text{inr } y \mapsto \text{right} \} \rightsquigarrow \text{left}(x) \\
\\
\text{match}(\text{inr } y) \{ \text{inl } x \mapsto \text{left}; \text{inr } y \mapsto \text{right} \} \rightsquigarrow \text{right}(y) \\
\\
\frac{\Gamma \vdash e : (X + Y)}{e \equiv \text{match}(e) \{ \text{inl } x \mapsto \text{inl } x; \text{inr } y \mapsto \text{inr } y \}}
\end{array}$$

Если у нас есть два типа данных X и Y , можно задать тип пар $X \times Y$.

$$\begin{array}{c}
\frac{X \quad Y}{X \times Y}(\text{pair}) \qquad \frac{\Gamma \vdash x : X \quad \Gamma \vdash y : Y}{\langle x, y \rangle : X \times Y} \qquad \frac{\Gamma \vdash p : X \times Y}{\Gamma \vdash p_{\text{fst}} : X} \qquad \frac{\Gamma \vdash p : X \times Y}{\Gamma \vdash p_{\text{snd}} : Y} \\
\\
\langle x, y \rangle_{\text{fst}} \rightsquigarrow x \qquad \langle x, y \rangle_{\text{snd}} \rightsquigarrow y \qquad \frac{\Gamma \vdash p : X \times Y}{p \equiv \langle p_{\text{fst}}, p_{\text{snd}} \rangle}
\end{array}$$

Если у нас есть два типа данных X и Y , можно задать тип функций $X \rightarrow Y$. Его каноническими элементами будут λ -термы — выражения типа Y нашего же языка с переменной типа X . Внутри этих выражений можно использовать в точности те синтаксические конструкции, которые мы ввели — оператор `if` и значения `tt` и `ff`, операторы `match` и конструкторы `inl`, `inr`, пары и проекции, ну и конечно же сами λ -термы и применение функций.

$$\begin{array}{c}
\frac{X \quad Y}{X \rightarrow Y}(\text{func}) \qquad \frac{\Gamma, x : X \vdash f : Y}{\Gamma \vdash \lambda(x : X) f : X \rightarrow Y} \qquad \frac{\Gamma \vdash f : X \rightarrow Y \quad \Gamma \vdash x : X}{\Gamma \vdash f x : Y} \\
\\
(\lambda(x : X) f(x)) x' \rightsquigarrow f(x') \qquad \frac{\Gamma \vdash f : X \rightarrow Y}{f \equiv \lambda(x : X) f x}
\end{array}$$

Получившийся язык программирования называется расширенным просто типизированным λ -исчислением, и все выразимые на нём типы данных суть конечные множества. Чтобы на нём были выразимы все конечные множества, нам нужно добавить ещё тип 1 с одним единственным значением и тип 0 вовсе без значений. Это в точности внутренний язык бидекартово замкнутых категорий.

(Обыкновенным) просто типизированным λ -исчислением называется подмножество данного языка, образуемое единичным типом, парами и функциями. Это — язык декартово замкнутых категорий. В отличие от расширенного варианта, в таком языке невозможно сконструировать типы с более чем одним значением, так что он тривиален как язык программирования, если в него не добавить дополнительные дополнительные базовые типы (такие как $\mathbb{B}$).

3.2. System T как язык программирования

Расширенное просто-типизированное λ -исчисление позволяет работать с конечными списками наперёд известного размера ($\text{Byte} := \mathbb{B}^8$), но не позволяет говорить о списках произвольного конечного размера

$$T^* = \mathbb{1} + T + T^2 + T^3 + \dots$$

Добавление типов списков произвольного размера как раз даст System T в максимальной формулировке — высокоуровневый язык программирования с минималистичной спецификацией. Во вводной лекции мы приводили все правила для типа списков T^* с оператором `reduce`; в этой лекции, для простоты изложения, мы вернемся к классическому гёделевскому варианту, где System T получается добавлением типов $\mathbb{B}$ и $\mathbb{N}$ (натуральные числа — частный случай списков $\mathbb{N} = \mathbb{1}^*$) и примитивной рекурсии к просто типизированному λ -исчислению (STLC), где имеются только функции и пары. Различия между этими формулировками System T тривиальны с формальной точки зрения.¹³

$$T ::= \mathbb{N} \mid \mathbb{B} \mid \mathbb{1} \mid T \rightarrow T \mid T \times T$$

По сравнению со STLC добавляются конструкторы и рекурсор для $\mathbb{N}$:

$$\frac{}{\mathbb{N} \text{ (nat)}} \quad \frac{}{\Gamma \vdash 0 : \mathbb{N}} \quad \frac{\Gamma \vdash n : \mathbb{N}}{\Gamma \vdash S n : \mathbb{N}} \quad \frac{\Gamma \vdash n : \mathbb{N} \quad \Gamma \vdash \text{step} : \mathbb{N} \rightarrow T \rightarrow T \quad \Gamma \vdash \text{base} : T}{\Gamma \vdash \text{iterate}(n) \text{ step base} : T}$$

$$\text{iterate}(0) \text{ step base} \rightsquigarrow \text{base} \quad \text{iterate}(S n) \text{ step base} \rightsquigarrow \text{step } n (\text{iterate}(n) \text{ step base})$$

Оператор `iterate` задаёт *примитивную рекурсию*: он разбирает число, имея на каждом шаге доступ к предшественнику и к результату рекурсивного вызова. Например, сложение записывается как

$$\text{add}(m)(n) = \text{iterate}(n)(\lambda k r. S r) m,$$

и вычисление `add(2)(2)` разворачивается в цепочку

$$S (\text{iterate}(1)(\dots) 2) \rightsquigarrow S S (\text{iterate}(0)(\dots) 2) \rightsquigarrow S S 2 = 4.$$

Классическое упражнение состоит в определении предшественника: прямой рекурсии «на единицу назад» у нас нет, но можно рекурсивно строить пару $(k-1, k)$ и в конце взять первую компоненту. Это типичный приём: результирующий тип T рекурсора может сам быть парой или функцией. Именно за счёт рекурсии на высших типах System T выразительно богаче примитивно-рекурсивных функций первого порядка: например, функция Аккермана записывается в нём без труда.

При этом язык тотален: все программы System T завершаются. Неформальный аргумент состоит в том, что каждый `iterate` съедает конкретное натуральное число, а типизированное λ -исчисление само по себе не содержит ни общей рекурсии, ни комбинаторов неподвижной точки. Самая мощная формулировка тотальности через отношение редукции называется сильной нормализацией и требует отдельного доказательства.¹⁴ Сконструированный ниже нормализатор реализует утверждение *для каждого терма существует нормальная форма*; конечность любой последовательности редукций из этого не следует.

¹³Типы $\mathbb{B}$ и $\mathbb{1}$, вообще говоря, тоже необязательны.

¹⁴Как правило, строящееся методом *логических отношений* (logical relations).

Внутренне типизированное (intrinsically typed) представление

В формализации синтаксис и типизация совмещены в одном индуктивном семействе:

```
Inductive tm : cxt -> ty -> Type :=
| tzero : tm Γ tN
| tsuc  : tm Γ tN -> tm Γ tN
| trec  : tm Γ T -> tm Γ (tN ⇒ T ⇒ T) -> tm Γ tN -> tm Γ T
| tvar  : var Γ T -> tm Γ T
| tlam  : tm (S :: Γ) T -> tm Γ (S ⇒ T)
| tapp  : tm Γ (S ⇒ T) -> tm Γ S -> tm Γ T
(* ... *)
```

Конструкторы `tm` непосредственно повторяют правила естественной дедукции: контекст представляет собой список типов, переменные представляют собой типизированные индексы де Брёйна (`var Γ T` свидетельствует, что T лежит в Γ), а `tlam` расширяет контекст. Некорректно типизированный терм таким способом не строится, поэтому отдельное отношение типизации и лемма о сохранении типов при редукции не нужны.

Открытые термы и заблокированные вычисления

Замкнутая программа типа $\mathbb{N}$ всегда вычисляется до числового значения.¹⁵ Интереснее ведут себя открытые термы. Сравним (при нашем определении сложения, где рекурсия ведётся по второму аргументу):

$$x + 2 \rightsquigarrow^* SSx, \quad 2 + x \rightsquigarrow^* \text{iterate}(x)(\lambda kr. Sr) 2.$$

Первый терм вычислился до конструкторов; второй *заблокирован*: итератору нечего разбирать, пока x неизвестен. Нормализация открытых термов представляет собой символьное исполнение, а не просто вычисление замкнутых данных. Заблокированные («нейтральные») подвыражения при этом сохраняются в нормальной форме вместе с уже вычисленными фрагментами.

3.3. Нормализация вычислением (NbE)

Нормализация нужна, чтобы упрощать и сравнивать программы System T, в том числе программы, которые ниже будут извлечены из доказательств интерпретацией «Диалектика». При упрощении нормализатор заранее вычисляет фрагменты, аргументы которых уже известны, и подставляет полученные значения в окружающий код. В процессе исчезают применения функций, проекции и шаги итерации, которые больше не зависят от входных данных, а в остаточной программе видно, какие неизвестные переменные блокируют дальнейшее вычисление.

Сравнение программ ставит другую задачу. Одна программа может быть записана до или после β -редукции, вычисления проекций и итератора либо η -развёртки, поэтому две её записи не обязаны посимвольно совпадать. Хотелось бы приводить такие записи к одной канонической форме и сравнивать уже результаты нормализации.

Дальше мы будем рассматривать упрощённый вариант System T, используемый в формализации. Напомним, что мы уже работаем в языке без типов сумм; кроме того, в дальнейшем мы не используем η -правила для базовых типов $\mathbb{N}$, $\mathbb{B}$ и $\mathbb{1}$. Свободная переменная базового типа не вычисляется и сохраняется в нормальной форме; переменные и заблокированные на них вычисления мы

¹⁵Это свойство называется *каноничностью*.

будем называть *нейтральными термами*, или коротко *нейтралями*. Другими словами, η -развёртка выполняется только для функций и пар, а приведённое в первом разделе правило уникальности для булевых значений не входит в набор преобразований, проверяемый нормализатором. Можно построить варианты NbE и для расширенных вариантов равенства, но у них будут другие нормальные формы и другой контракт алгоритма.

Будем писать $t \equiv u$ и говорить, что термы *равны по определению* (definitionally equal),¹⁶ если один можно преобразовать в другой с помощью β -правила для функций, правил вычисления проекций, уравнений `iterate` и `if`, η -правил для функций и пар, а также обратных преобразований и замен внутри составных термов.

Почему бы не сравнивать программы непосредственно по их поведению? Для языка с общей рекурсией точная семантическая эквивалентность программ неразрешима. Равенство по определению служит разрешимым компромиссом: в System T каждый терм имеет нормальную форму, и нормализатор даёт

$$t \equiv u \iff \text{norm}(t) = \text{norm}(u).$$

Справа стоит обычное синтаксическое совпадение нормальных форм. Значит, равенство по определению разрешается сравнением нормальных форм. Так нормализация используется для проверки равенства по определению в зависимо типизированных языках, где типы могут содержать программы (об этом подробно в лекции о зависимых типах).

Заметим, что таким образом решается только выбранное *равенство по определению*, а не полная эквивалентность. Например, рассмотренные выше термы $x + 2$ и $2 + x$ равны по теореме о коммутативности сложения, доказательство которой требует индукции. Однако их нормальные формы различны: $\text{SS } x$ и $\text{iterate}(x)(\lambda k r. S r) 2$. Нормальная форма также не обязана сохранять структуру стоимости и общих подвыражений: β -редукция и η -развёртка могут дублировать синтаксис и уничтожать информацию о переиспользуемых подвыражениях.

Проговорим еще раз, что именно реализует NbE. Тотальность `norm` даёт каждому типизированному терму вычислимую каноническую форму, а приведённый выше критерий позволяет разрешать равенство по определению. Нормализатор называется *полным* (full), поскольку нормализует в том числе под λ , в отличие от вычислителя, останавливающегося на слабой головной нормальной форме (WHNF), то есть вычисляющего до первого внешнего конструктора или λ и не заходящего внутрь.

Сильная нормализация утверждает больше: любая последовательность редукций должна быть конечной, независимо от того, какое применимое правило и в каком месте терма выбирают на каждом шаге. Тотальный алгоритм `norm` сам по себе этого не доказывает. В формализации нет малшагового (small-step) отношения редукции, о последовательностях которого можно было бы сформулировать такую теорему. Вместо этого для функции `norm` доказаны здравость (soundness) и полнота (completeness), сформулированные ниже в «Контракте алгоритма».

¹⁶Также говорят о *конверсии* между такими термами.

Спецификация выхода: нормальные и нейтральные формы

Выход нормализатора описывается двумя взаимно определёнными грамматиками:

$v ::= 0 \mid S v \mid \mathbf{tt} \mid \mathbf{ff} \mid () \mid \lambda x. v \mid \langle v, v \rangle \mid u$ (u входит в v только при базовом типе)

$u ::= x \mid u v \mid \mathbf{iterate}(u) v v \mid \mathbf{if}(u) v v \mid u_{\text{fst}} \mid u_{\text{snd}}$

Нормальная форма v не содержит применений вида $(\lambda x. t) u$, к которым можно применить β -правило; все известные рекурсоры и условные конструкции вычислены, а заблокированные элиминации сохранены явно в виде нейтральных термов u , то есть цепочек элиминаций, упирающихся в переменную. Нейтрали входят в нормальные формы *только при базовом типе*: нейтраль функционального типа η -разворачивается до λ , а парного до пары. В полученных η -длинных формах каждая функция записана как λ -абстракция, а каждая пара как $\langle \dots, \dots \rangle$; такие формы служат каноническими представителями $\beta\eta$ -равенства.

Идея алгоритма

Для получения нормальной формы можно последовательно применять β -правило, правила проекций и уравнения `iterate` и `if` во всём терме, включая подтермы под λ , пока ни одно правило больше не применимо. Такая «сильная редукция» в духе [Grégoire and Leroy \(2002\)](#) работает, но каждый β -шаг требует обхода тела функции, сдвига индексов де Брёйна или переименования связанных переменных и иногда копирования подставляемого аргумента. Повторные обходы усложняют реализацию и создают промежуточный синтаксис. Кроме того, этот способ не выполняет η -развёртку: нейтральную функцию f он оставит как f , тогда как каноническим представителем при $\beta\eta$ -равенстве является η -длинная форма $\lambda x. f x$. NbE использует другую схему: «сначала проинтерпретируй, потом прочитай обратно».

синтаксис $\xrightarrow{\text{eval}}$ семантическое значение $\xrightarrow{\text{reify}}$ каноническая форма

Идея его восходит к семантическому доказательству Мартин-Лёфа о том, что каждый терм его предикативной интуиционистской теории зависимых типов имеет нормальную форму ([Martin-Löf, 1975](#)). Бергер и Швихтенберг превратили её в самостоятельный алгоритм под названием «обращение функционала вычисления» ([Berger and Schwichtenberg, 1991](#)). Позже Данви переоткрыл ту же схему как *type-directed partial evaluation* (TDPE) ([Danvy, 1996](#)), а Катарина Кокан объяснила её через связь синтаксиса с его моделью ([Coquand, 2002](#)).

Каждому типу объектного языка сопоставляется его семантическая интерпретация в метаязыке, в нашем случае в самом Rocq (либо OCaml, если рассматривать экстрагированную в него имплементацию):

$\llbracket B \rrbracket_\Gamma = \text{nf } \Gamma \ B$ (нормальные формы базового типа B в контексте Γ)

$\llbracket A \rightarrow B \rrbracket_\Gamma = \forall \Delta. \text{ope } \Delta \ \Gamma \rightarrow \llbracket A \rrbracket_\Delta \rightarrow \llbracket B \rrbracket_\Delta$ (метаязыковые функции)

$\llbracket A \times B \rrbracket_\Gamma = \llbracket A \rrbracket_\Gamma \times \llbracket B \rrbracket_\Gamma$ (метаязыковые пары)

Окружение $\rho : \text{Env } \Delta \ \Gamma$ сопоставляет каждой переменной $x : T$ из контекста Γ семантическое значение из $\llbracket T \rrbracket_\Delta$. Два индекса играют разные роли: Γ представляет исходный контекст терма, переменные которого нужно оценить, а Δ представляет

контекст, в котором живут полученные семантические значения. В тождественном окружении эти контексты совпадают. Различать их необходимо при переходе в расширенный контекст, например под λ . Вычислитель eval берёт значения свободных переменных из окружения, а при вычислении тела λ расширяет окружение семантическим значением аргумента.

Важно различать два обозначения: $\llbracket T \rrbracket_\Delta$ задаёт семантическую область для типа T , а eval интерпретирует *терм* в окружении и возвращает элемент этой области. Формально,

$$\text{eval} : \text{tm } \Gamma \ T \rightarrow \text{Env } \Delta \ \Gamma \rightarrow \llbracket T \rrbracket_\Delta.$$

Таким образом, $\text{eval}(t)(\rho)$ представляет собой семантическое значение терма t при значениях свободных переменных из ρ . Если бы для термов использовалась запись $\llbracket t \rrbracket_\rho$, она обозначала бы то же самое; ниже вместо неё везде пишется $\text{eval}(t)(\rho)$.

Вместе эти определения задают денотационную семантику: типы переходят в семантические области, контексты в окружения, а термы в значения, зависящие от окружения. В случае функций значение дополнительно индексировано контекстом и квантифицировано по оре $\Delta \ \Gamma$. ОРЕ (order-preserving embedding) представляет собой конструктивное доказательство того, что Γ вкладывается в Δ с сохранением порядка (другими словами, Δ получен из Γ только добавлением переменных).

Если для наглядности обозначить позиции контекста буквами, вложение может выглядеть так:

$$\begin{array}{ccccc} \Gamma : [a & & c & & e] \\ \downarrow & & \downarrow & & \downarrow \\ \Delta : [a & b & c & d & e] \end{array}$$

ОРЕ сохраняет порядок старых позиций a, c, e и пропускает добавленные позиции b, d .

Другими словами, мы требуем, чтобы семантическая функция работала и во всех расширениях своего контекста. Это нужно при чтении под λ : там появляется свежая переменная, контекст расширяется и уже вычисленные значения требуется перенести в новый контекст. В терминах теории категорий такие зависящие от контекста значения образуют предпучки на категории контекстов.

Обратное чтение задаётся парой функций, определяемых взаимной рекурсией *по типу*:

- $\text{reify}_T : \llbracket T \rrbracket_\Gamma \rightarrow \text{nf } \Gamma \ T$ превращает семантическое значение в каноническую нормальную форму;
- $\text{reflect}_T : \text{ne } \Gamma \ T \rightarrow \llbracket T \rrbracket_\Gamma$ наделяет нейтральный терм (например, переменную) семантическим значением.

Для базового типа обе функции являются тождественными вложениями. Для функций определяем:

$$\text{reify}_{A \rightarrow B}(f) = \lambda x. \text{reify}_B(f(\text{reflect}_A(x))), \quad \text{reflect}_{A \rightarrow B}(u)(a) = \text{reflect}_B(u(\text{reify}_A(a))).$$

Чтобы прочитать семантическую функцию, мы вводим свежую переменную x , *отражаем* её в семантическое значение, применяем функцию и рекурсивно читаем результат. Так алгоритм нормализует под λ и выполняет η -развёртку. Отражение нейтральной функции действует в обратную сторону: её применение

остаётся нейтральным термом, а семантический аргумент сначала читается обратно в синтаксис.

По устройству `eval` похож на интерпретатор больших шагов (big-step) с окружениями, но возвращает не обычные операционные значения, а значения семантического домена NbE: функции и пары метаязыка, а для базовых типов нормальные и нейтральные формы. Чтобы получить синтаксический результат, после `eval` применяется `reify`. При вычислении λ переходит в метаязыковую функцию, применение в применение, а для итератора `eval` сначала вычисляет его натуральный аргумент. Поскольку $\llbracket \mathbb{N} \rrbracket_{\Gamma} = \text{nf } \Gamma \ \mathbb{N}$, полученное значение имеет вид $0, S \ n$ либо является нейтральным. В первых двух случаях вычислитель применяет соответствующее уравнение `iterate`. В последнем он строит нейтральный `iterate`, у которого шаг и база уже прочитаны обратно.

Нормализатор задаётся композицией `reify` и `eval`:

$$\text{norm}(t) = \text{reify}(\text{eval}(t)(\rho_{\text{id}})), \quad \rho_{\text{id}}(x) = \text{reflect}(x).$$

Терм вычисляется в окружении, где каждая свободная переменная отражена в саму себя. Подстановку и β -редукцию выполняет метаязык, а синтаксис строится заново при обратном чтении. Тотальность конструкции опирается на тотальность System T: с общей рекурсией `eval` мог бы зависнуть, в том числе внутри `reify` при применении семантической функции к свежей переменной.

Примеры из формализации (`Tests.v`):

- замкнутое вычисление: $\text{norm}(2 + 2) = 4$;
- η -развёртка: $\text{norm}(f) = \lambda x. f x$ при свободной $f : \mathbb{N} \rightarrow \mathbb{N}$, и $\text{norm}(p) = \langle p_{\text{fst}}, p_{\text{snd}} \rangle$ при свободной $p : A \times B$;
- заблокированная рекурсия: $\text{norm}((1+1)+x) = \text{iterate}(x)(\lambda k r. S \ r) \ 2$; известная часть вычислена, нейтральный остаток сохранён.

Контракт алгоритма

Метатеория формализации доказывает следующие утверждения (мы принимаем их без доказательства; используется техника логических отношений, о которой пойдёт речь в следующих лекциях):

- *здравость*: $\text{norm}(t) \equiv t$, то есть нормализация не меняет программу с точностью до равенства по определению;
- *полнота*: если $t \equiv u$, то $\text{norm}(t)$ и $\text{norm}(u)$ синтаксически совпадают.

Следовательно, $t \equiv u \iff \text{norm}(t) = \text{norm}(u)$: равенство по определению разрешимо, нормальные формы единственны, а каждый замкнутый терм базового типа равен конструкторному значению.

Схема «вычислить в семантике, затем прочесть обратно» применяется и за пределами этого нормализатора. Вслед за перечнем Линдли эти применения можно сгруппировать в две области (Lindley, 2016). В программировании это встраиваемые DSL, частичные вычисления и эффективная β -редукция с помощью замыканий метаязыка. В исследованиях по основаниям логики и языков это исполнимая денотационная семантика, нормализация доказательств и проверка равенства по определению в теориях типов. Реализации NbE используются в тайпчекерах зависимо типизированных языков и в Isabelle/HOL, где команда `value [nbe]` вычисляет термы методом NbE. Ещё один пример даёт конфигурационный язык Dhall. Работы Киселёва с соавторами показывают, как частичное вычисление и схема, вдохновлённая NbE, применяются для оптимизаций встраиваемых DSL (Carette et al., 2009; Kiselyov, 2014).

3.4. Реализуемость и интерпретация «Диалектика»

Исторически Диалектика возникла из программы Гильберта. Она требовала доказать непротиворечивость классической арифметики Пеано (РА) «финитными» средствами: элементарными рассуждениями о числах, формулах и доказательствах как конечных последовательностях символов, без обращения к завершённой бесконечности. Надо сказать, точной формальной границы этого понятия Гильберт не задал.

После теорем Гёделя о неполноте эту финитную точку зрения пришлось расширить. Гёдель предложил считать конструктивными объектами не только числовые вычисления, но и *примитивно-рекурсивные функционалы конечных типов*, образующие System T. Здесь «конечный» описывает построение типа из $\mathbb{N}$ с помощью стрелок, а не число его значений: $\mathbb{N} \rightarrow \mathbb{N}$ тоже является конечным типом. В 1941 году Гёдель изложил функциональную интерпретацию арифметики Гейтинга (НА), а в 1958 году опубликовал её в журнале *Dialectica* (Gödel, 1958) (откуда и название). НА имеет тот же язык, арифметические аксиомы и индукцию, что и РА, но использует интуиционистскую, а не классическую логику; добавление закона исключённого третьего превращает НА в РА. Поэтому классические доказательства РА сначала переводятся в НА с помощью перевода двойного отрицания,¹⁷ а затем интерпретируются Диалектикой. Извлечение программ из доказательств представляет собой более позднее прочтение той же конструкции; современный обзор дан в Avigad and Feferman (1998).

Ниже рассматривается упрощённый вариант интерпретации, использованный для формализации. Формулы в нём строятся из связок $\wedge, \vee, \supset$, кванторов по натуральным числам и разрешимых атомов. Каждый атом уже задан булевой программой System T; так можно закодировать равенство, порядок и другие примитивно-рекурсивные предикаты. Это позволяет сосредоточиться на том, как Диалектика преобразует логическую структуру формулы.

В более буквальной формализации НА можно было бы "честно" задать синтаксис арифметических термов и атомарных формул, а затем дополнительно компилировать их в булевы программы System T. Такой вариант добавил бы рутинный слой трансляции и технические леммы о подстановке и переводе выводов, но не изменил бы самой сути компилятора Диалектики.

Свидетели против контрпримеров

Каждой формуле A интерпретация сопоставляет два типа System T и бескванторную матрицу:

$$A \rightsquigarrow \exists w : W(A). \forall c : C(A). |A|(w, c),$$

где $|A|(w, c)$ представляет собой разрешимый тест. Удобно читать это как игру: $W(A)$ задаёт тип *свидетелей* (witness, ходов защитника), $C(A)$ задаёт тип *контрпримеров* (counter, ходов оппонента), и свидетель w «побеждает» контрпример c , если $|A|(w, c)$ истинно. Типы определяются рекурсией по формуле:

A	$W(A)$	$C(A)$
атом	$\mathbb{1}$	$\mathbb{1}$
$X \wedge Y$	$W(X) \times W(Y)$	$C(X) \times C(Y)$
$X \vee Y$	$\mathbb{B} \times (W(X) \times W(Y))$	$C(X) \times C(Y)$
$X \supset Y$	$(W(X) \rightarrow W(Y)) \times (W(X) \rightarrow C(Y) \rightarrow C(X))$	$W(X) \times C(Y)$
$\forall x. X$	$\mathbb{N} \rightarrow W(X)$	$\mathbb{N} \times C(X)$
$\exists x. X$	$\mathbb{N} \times W(X)$	$C(X)$

¹⁷Он систематически добавляет двойные отрицания и тем самым переводит классические доказательства в интуиционистские.

Единичные типы в строке для атома не означают, что атом автоматически истинен. Они означают только, что для него не извлекаются отдельные свидетель и контрпример: вся содержательная проверка уже находится в самом булевом атоме. Если атом задан программой $p : \mathbb{B}$, то

$$|p|((\cdot), (\cdot)) = p.$$

Для конъюнкции $X \wedge Y$ свидетель $\langle w_X, w_Y \rangle$ содержит свидетелей обеих частей, а контрпример $\langle c_X, c_Y \rangle$ содержит по вызову для каждой из них. Матрица требует, чтобы w_X победил c_X , а w_Y победил c_Y .

Для дизъюнкции $X \vee Y$ свидетель имеет вид $\langle b, \langle w_X, w_Y \rangle \rangle$. Булев флаг b выбирает защищаемую часть: tt выбирает X , а ff выбирает Y . Тип хранит кандидатов для обеих частей, потому что в этой кодировке нет типа суммы; неиспользуемый компонент можно заполнить стандартным значением соответствующего типа. Каждый тип рассматриваемого варианта System T обитаем, поэтому такой заполнитель всегда существует. Контрпример $\langle c_X, c_Y \rangle$ предлагает вызов для каждой части, а матрица по флагу выбирает только c_X или только c_Y .

Свидетель импликации $X \supset Y$ представляет собой *пару функций* $\langle f, f^\leftarrow \rangle$. Обе функции входят в свидетель: прямая компонента f превращает свидетель X в свидетель Y , а обратная компонента $f^\leftarrow$ превращает контрпример к Y в контрпример к X , используя свидетель X . Сама обратная компонента контрпримером не является. Контрпример к импликации имеет вид $\langle w_X, c_Y \rangle$: он предъявляет свидетель посылки и вызов заключению, который $f^\leftarrow$ переносит назад к посылке. Та же структура прямого и обратного прохода встречается в обратном дифференцировании и связана с линейной логикой.

Для читателя, знакомого с линзами из функционального программирования, эта пара имеет тип *биморфной линзы* из $(W(X), C(X))$ в $(W(Y), C(Y))$. Линза объединяет прямое отображение $f : W(X) \rightarrow W(Y)$ с обратным отображением $f^\leftarrow : W(Y) \rightarrow C(X)$ (Hedges, 2018). Правило композиции импликаций компилируется в композицию таких линз:

$$\langle g, g^\leftarrow \rangle \circ \langle f, f^\leftarrow \rangle = \langle \lambda w. g(f(w)), \lambda w c. f^\leftarrow(w)(g^\leftarrow(f(w))(c)) \rangle.$$

Здесь используется только тип и композиция линз; обычные законы линз не утверждаются.

Типы кванторов также заранее задают форму программы. Например,

$$W(\forall x. \exists y. y = Sx) = \mathbb{N} \rightarrow (\mathbb{N} \times \mathbb{1})$$

Поэтому свидетель этой теоремы обязан быть программой, вычисляющей по x значение y ; тип уже предсказывает форму извлечённой программы.

Арифметические термы внутри формул уже являются термами System T, поэтому трансляция меняет логическую структуру, но не арифметику. Матрица $|A|$ компилируется структурной рекурсией в программу System T

$$\text{diaT } A : W(A) \rightarrow C(A) \rightarrow \mathbb{B},$$

например $|X \wedge Y|(\langle w_X, w_Y \rangle, \langle c_X, c_Y \rangle) = |X|(w_X, c_X) \ \&\& \ |Y|(w_Y, c_Y)$, а $|\exists x. X|(\langle n, w \rangle, c) = |X[n/x]|(w, c)$. Матрица исполняется во время вычислений: извлечённые реализаторы вызывают её для выбора контрпримеров.

Компиляция доказательств

До сих пор мы описывали только трансляцию формул: формуле A сопоставлялись типы $W(A)$ и $C(A)$ и матрица $|A|$. Чтобы извлечь программу, нужно также перевести доказательство этой формулы.

В описанной раньше натуральной дедукции суждение $\Gamma \vdash A$ содержит явный список логических предпосылок Γ , которыми правила оперируют непосредственно. В гильбертовском исчислении списка предпосылок нет: доказательство собирается из схем аксиом с помощью *modus ponens* и нескольких других правил. Поэтому в типе `proof` $\Gamma \vdash A$, кодирующем гильбертов вывод, тот же символ Γ обозначает только контекст свободных числовых переменных, а не логические гипотезы.

Для программиста гильбертовское доказательство похоже на программу для стекового языка: схема аксиом кладёт готовое доказательство на стек, а *modus ponens* забирает доказательства $A \supset B$ и A и возвращает доказательство B . Натуральная дедукция, напротив, ближе к λ -исчислению: предпосылки играют роль переменных, а введение импликации соответствует λ -абстракции.

Оба представления могут задавать одну и ту же логику. Гильбертовский вариант выбран здесь ради более простой формализации компилятора: каждой схеме аксиом и каждому правилу соответствует отдельный случай трансляции, а для списков гипотез не приходится доказывать технические леммы о перестановке предпосылок и подстановке одного вывода в другой.

Интерпретация доказательств задаётся компилятором

$$\text{wit} : \text{proof } \Gamma \vdash A \rightarrow \text{tm } \Gamma (W(A))$$

Компилятор `wit` строит только свидетель типа $W(A)$. Контрпример типа $C(A)$ он не извлекает: тот поступает извне как вызов, с которым свидетель проверяется матрицей $|A|$. В частности, обратные компоненты для импликаций входят в построенный свидетель и преобразуют контрпримеры, но не являются отдельным результатом `wit`.

Компилятор определяется рекурсией по выводу гильбертовского исчисления НА с индукцией. Внутренняя типизация гарантирует, что результат `wit` имеет тип $W(A)$. Например:

- *modus ponens*: из свидетеля u импликации берётся его прямая компонента и применяется к свидетелю посылки, $\text{wit} = u_{\text{fst}}(a)$;
- *введение* $\exists$: терм t , использованный в правиле, упаковывается в пару $\langle t, w \rangle$. Свидетель существования не «синтезируется» интерпретацией, а извлекается из доказательства, где он уже содержался;
- *сжатие* (contraction) $P \supset P \wedge P$: прямая компонента свидетеля импликации дублирует свидетель P , а обратная компонента получает два контрпримера и должна вернуть один. Извлечённая программа запускает матрицу и выбирает тот контрпример, который действительно опровергает:

$$\langle \lambda w. \langle w, w \rangle, \lambda w \langle c_1, c_2 \rangle. \text{if}(|P|(w, c_1)) c_2 c_1 \rangle.$$

Так логическая структура превращается в поток управления времени исполнения; выбор возможен потому, что матрица разрешима.

Правило индукции представлено аксиомой

$$(Q(0) \wedge \forall k. (Q(k) \supset Q(Sk))) \supset \forall n. Q(n).$$

Её свидетель состоит из пары примитивных рекурсий. Прямая компонента строит свидетеля $Q(n)$ снизу вверх: w_0 служит базовым свидетелем, s_k служит

свидетелем шага $Q(k) \supset Q(Sk)$, а $w_{k+1} = (s_k)_{\text{fst}}(w_k)$. Обратная компонента того же свидетеля преобразует контрпримеры, выполняя гёделевский *поиск контрпримера*: получив вызов против $Q(n)$, она протаскивает его назад сквозь свидетелей шагов, на каждой ступени запуская матрицу $|Q(k)|$, и возвращает первую ступень, которую можно «обвинить», то есть контрпример базе или контрпример конкретному шагу. Каждая компонента использует по одному *iterate*, поэтому индукция и при исполнении остаётся примитивной рекурсией.

Заметим, что интерпретация дополнительно реализует два интуиционистски *не выводимых* принципа. Здесь $\neg A$ является сокращением для $A \supset \text{False}$, где False есть атом, заданный булевой программой ff . Бескванторный принцип Маркова и независимость посылки для $\forall$ -посылок имеют вид

$$\begin{aligned} \text{MP} : & (\neg \forall x. \neg P(x)) \supset \exists x. P(x), \\ \text{IP}_{\forall} : & ((\forall x. P(x)) \supset \exists y. Q(y)) \supset \exists y. ((\forall x. P(x)) \supset Q(y)). \end{aligned}$$

В MP предикат P берётся из предусмотренного формализацией разрешимого бескванторного фрагмента. В формализации оба принципа добавлены как аксиомы с реализаторами, так что Диалектика «конструктивизирует» и эти расширения интуиционистской логики.

Чтобы сформулировать валидность извлечённого свидетеля, мы строим более простую модель System T, чем модель NbE. Она интерпретирует $\mathbb{N}$, $\mathbb{B}$ и $\mathbb{1}$ обычными типами значений метаязыка, а функции и пары метаязыковыми функциями и парами. Здесь не нужны нейтральные значения, OPE и обратное чтение, поскольку модель не строит нормальную форму, а только вычисляет значение терма.

Окружение $\rho : \text{cxt} \text{den } \Gamma$ этой модели сопоставляет свободным переменным числа, булевы значения, функции и пары, а $\text{tmden}(t, \rho)$ вычисляет значение терма t . Относительно этой модели предикат $\text{valid}(A, t)$ означает, что программа $t : W(A)$ при любых значениях свободных переменных побеждает любой контрпример. Функция dia вычисляет матрицу, поэтому содержательно определение имеет вид

$$\text{valid}(A, t) \iff \forall \rho. \forall c. \text{dia}(A, \rho, \text{tmden}(t, \rho), c) = \text{tt}.$$

Общая теорема здравости *soundness* утверждает, что для всякого вывода d формулы A реализатор $\text{wit}(d)$ побеждает любой контрпример в смысле предиката valid .

Полнота здесь не утверждается: из существования программы $t : W(A)$, для которой выполнено $\text{valid}(A, t)$, в общем случае нельзя восстановить доказательство формулы A в НА. Диалектика служит интерпретацией доказательств и методом извлечения программ, а не критерием доказуемости.

Канонические реализаторы

Два алгоритма (NbE и Диалектика) образуют цепочку:

$$\text{вывод } d \text{ формулы } A \xrightarrow{\text{wit}} \text{программа } \text{wit}(d) : W(A) \xrightarrow{\text{norm}} \text{каноническая программа}$$

Первая стрелка реализует Гёделевскую идею: программу $\text{wit}(d)$ уже можно исполнять. Отдельный этап *norm* приводит её к канонической форме. Алгоритмы нормализации программ и трансляции доказательств возникли из разных задач, но здесь работают с одним синтаксисом System T.

Композиция двух алгоритмов записывается тривиально:

Definition $\text{realizer } \{\Gamma\} \{A : \text{prp } \Gamma\} (d : \text{proof } \Gamma A) : \text{nf } \Gamma (W A) :=$
 $\text{norm } (\text{wit } d).$

Заметим также, что для замкнутой формулы окружение в определении `valid` отсутствует, а здравость можно сформулировать синтаксически: $\text{norm}(\text{diaT } A \cdot \text{wit}(d) \cdot c) = \text{tt}$ для любого замкнутого c .

До запуска NbE компилятор `wit` выдаёт обычные исполнимые термы System T. Нормализация устраняет в них промежуточные конструкции и строит каноническую $\beta\eta$ -длинную форму. Для свидетелей высших типов такая форма показывает функциональную структуру, которую нельзя свести к одному числовому ответу. Атомарные равенства в типах свидетелей стираются до единичного типа $\mathbb{1}$ и представлены значением `()`; само равенство при этом проверяется булевой матрицей. Примеры из формализации (`Examples.v`; через r_R обозначен фиксированный реализатор правила R гильбертовского исчисления):

- $\text{realizer}(\vdash \exists y. y = 2) = \langle 2, () \rangle$;
- для доказательства ex_succ формулы $\vdash \forall x. \exists y. y = S x$ компилятор сначала выдаёт

$$\text{wit}(\text{ex_succ}) = r_{\text{mp}}(r_{\forall I}(r_{\text{mp}}(r_{\text{cur}}(r_{\wedge E_1}), \\ r_{\text{mp}}(r_{\exists I}(S x), ())), \\ ()).$$

Композиция правил оставляет здесь вспомогательные применения λ -функций. После NbE получается $\text{realizer}(\text{ex_succ}) = \lambda x. \langle S x, () \rangle$: по x программа вычисляет свидетеля $y = S x$;

- для аксиомы сжатия $P \supset P \wedge P$ стратегия заметна уже в результате компиляции:

$$\text{wit}(\text{ax_and_contr}) = r_{\wedge C} \\ = \langle \lambda w. \langle w, w \rangle, \\ \lambda w \langle c_1, c_2 \rangle. \\ \text{if}(\text{diaT } P \cdot w \cdot c_1) c_2 c_1 \rangle.$$

Здесь NbE не меняет содержательную структуру программы, а лишь раскрывает скомпилированную матрицу `diaT P` и приводит терм к канонической паре

$$\langle \lambda w. \langle w, w \rangle, \lambda w \langle c_1, c_2 \rangle. \text{if}(|P|(w, c_1)) c_2 c_1 \rangle,$$

где первая функция дублирует свидетель, а вторая запускает матрицу $|P|$ и направляет назад тот контрпример, который может опровергнуть P ;

- доказательство $\vdash \forall x. 0 + x = x$ индукцией даёт $\lambda n. \text{iterate}(n)(\lambda _ . ())()$, скелет индукции со стёртым до единиц содержанием. При определении сложения с рекурсией по второму аргументу это не равенство по определению. Поскольку равенство является атомом, извлечённый свидетель возвращает только `()`. Арифметическое условие $0 + x = x$ остаётся в булевой матрице, а `iterate` сохраняет структуру доказательства индукцией;
- доказательство $\vdash \exists y. y = 0$ через *принцип Маркова* (нарочито неконструктивное: из $\neg \forall y. \neg(y = 0)$) всё равно нормализуется в $\langle 0, () \rangle$. Нормализация прогоняет контрпримерную машинерию негативного доказательства и добывает конкретного свидетеля.

Итак, System T служит тотальным целевым языком, Дialeктика компилирует в него логические выводы, а NbE при необходимости приводит полученные программы к каноническому виду.

У трансляции Гёделя есть варианты Диллера-Нама, Крайзеля и Стейна, а также монотонная интерпретация Коленбаха. Последняя лежит в основе `proof`

mining, то есть систематического извлечения количественных оценок из неконструктивных доказательств в математическом анализе (Kohlenbach, 2008).