

Internalizing Extensions in Lattices of Type Theories

Research Qualifier

Jonathan Chan

21 November 2024

Abstract

Many proof assistants allow the use of features and axioms that increase their expressive power. However, these extensions must be used with care, as some combinations are known to lead to logical inconsistencies. Therefore, proof assistants include mechanisms that track which extensions are used in a proof development or module, ensuring that incompatible extensions are not used simultaneously.

Unfortunately, existing extension tracking mechanisms are external to the type system. This means that we cannot specify precisely which extensions a definition depends on. Having the ability to write more precise specifications means we are not picking an overapproximation of the extensions needed, which prevents reusing definitions in the presence of incompatible extensions. Furthermore, we cannot refer to definitions that use incompatible extensions even if they are never used in inconsistent ways. The reasoning principles of one extension therefore cannot be used as a metatheory to reason about the properties of an incompatible extension.

In this report, I explore the use of the Dependent Calculus of Indistinguishability (DCOI) [Liu et al., 2024b] for extension tracking. DCOI is a dependent type system with dependency tracking, where terms and variables are assigned dependency levels alongside their types. These dependency levels form a lattice that describes which levels are permitted to access what. To instead track extensions, each set of extensions would correspond to a dependency level, and the lattice would describe how extensions are permitted to interact.

1 Introduction

At the core of a proof assistant founded on the Curry–Howard correspondence is a type checker that validates a proof of a proposition—a term inhabiting a corresponding type in a dependent type theory. These type theories are typically based on some flavour of Martin-Löf Type Theory (MLTT) [Martin-Löf, 1972] or the Calculus of Inductive Constructions (CIC) [Pfenning and Paulin-Mohring, 1990]. In practice, a proof assistant does not implement merely one type theory, but a whole host of them, as they include language extensions that augment or modify their reasoning power.

These type theoretic extensions consist of new typing rules, axioms, constructs, and/or definitional equalities. Because each extension embodies semantically distinct reasoning principles, enabling an extension results in a separate theory altogether. Unfortunately, some extensions are mutually incompatible and cannot be enabled simultaneously because

together they violate logical consistency. For instance, the *uniqueness of identity proofs*, which propositionally equates all proofs of the same equality, is incompatible with *univalence*, which adds provably distinct proofs of equality. In [Section 2](#), I describe these extensions and a few more found in select proof assistants, along with the ways some combinations are incompatible.

Proof assistants are careful to track the usage of language extensions to rule out inconsistencies. However, the tracking done by their type checkers is *external* to the type system: within the language itself, one cannot assert that a definition is permitted to depend on a particular extension, nor that it is prohibited from using an extension. Being able to specify exactly where the need for various extensions comes from helps us be more mindful about enabling them, especially if doing so prevents future reuse of definitions due to incompatible extensions. Furthermore, given two incompatible extensions, a definition in one extension’s theory may not be used at all in the another extension’s theory, not even just in a type. This means one theory cannot be used as a metatheory to prove properties about definitions in the other theory if those theories are incompatible. For instance, considering classical axioms as a language extension, one would not be able to explore what is constructively proveable about theories that use classical principles.

Ultimately, what is missing is a framework for describing fine-grained control of proofs and programs across multiple type theories, where even incompatible theories can interact in interesting ways. This fine-grained control of terms is reminiscent of type systems with *dependency analysis*. In such systems, terms are stratified by *dependency levels*, which can be thought of intuitively as permission levels tracking where terms may be used. Even if we do not have access to a particular level, terms at that level can still be manipulated and reasoned about opaquely as long as they are not inspected or evaluated.

I have worked on a dependent type system with dependency tracking, the Dependent Calculus of Indistinguishability (DCOI) [[Liu et al., 2024b](#)], along with its variant DCOI^ω [[Liu et al., 2025](#)], which is a logically consistent type theory. DCOI could potentially be used as a basis for a framework that internalizes extension tracking by stratifying the extensions’ corresponding type theories into hierarchies of dependency levels, where compatibility between extensions maps to the intuitive notion of permission. I describe the key properties of DCOI relevant for this application in [Section 3](#). I then speculate on the details of this mapping, lay out the objectives for such a framework, and list possible first steps towards accomplishing them in [Section 4](#). There is much prior and related work that this project relies on and relates to, which I divide into work I have personally contributed to ([Section 5](#)) and other work in this space ([Section 6](#)).

2 Proof assistant extensions in practice

To look at extensions in practice, let us focus on three popular proof assistants: Rocq [[Coq Development Team, 2022](#)], Agda [[Norell, 2007](#)], and Lean [[de Moura et al., 2015](#)]. Broadly speaking, they are all based on variants of MLTT or CIC, and have dependent functions, type universes, and inductive types.

2.1 Built-in features

Each of these proof assistants includes features that extend the power of their foundations; below are a few notable extensions, some of which are controlled by option flags.

Impredicativity. Rocq and Lean, being based on CIC, feature a universe `Prop` of propositions. This universe is *impredicative*, meaning that a universal quantification (*i.e.* dependent function type) $\forall(x : A). B$ is a proposition if B is a proposition, regardless of the universe in which A lives, which may be larger than `Prop`. Inductive types may also be defined in `Prop`, which permits its constructors to have argument types in universes larger than `Prop`. Such inductives are said to be *large*. An impredicative `Prop` is a default feature of Rocq and Lean, not an option that can be turned off.

Impredicativity allows for self-referential propositions by quantifying over `Prop`. For instance, given a proof that all propositions imply their double negation,

$$dn : \forall(P : \text{Prop}). P \rightarrow \neg\neg P,$$

the double negation of this proposition itself holds as well by self-application.

$$dn (\forall(P : \text{Prop}). P \rightarrow \neg\neg P) \quad dn : \neg\neg(\forall(P : \text{Prop}). P \rightarrow \neg\neg P)$$

In contrast, in the predicative setting, a quantification over a universe `Type0 : Type1` has type `Type1`, and in particular,

$$\Pi(A : \text{Type}_0). A \rightarrow \neg\neg A : \text{Type}_1,$$

so an element of this type may not be applied to the type itself.

Dually to universal quantification, we can define an existential quantification $\exists(x : A). B$ over a type A in a larger universe as a large inductive proposition, *i.e.* a dependent pair type in `Prop`. This enables us to state, for example, the surjectivity of a function over the naturals as a predicate, even though the naturals are not a proposition.

$$\begin{aligned} surj &: (\text{Nat} \rightarrow \text{Nat}) \rightarrow \text{Prop} \\ surj &:= \lambda f. \forall(y : \text{Nat}). \exists(x : \text{Nat}). f \, x \equiv y \end{aligned}$$

Definitional proof irrelevance. A universe of propositions is said to be *strict* when the inhabitants of its propositions are definitionally equal (*i.e.* proof irrelevant) and thus treated as interchangeable during type checking. Lean’s `Prop` is always strict, while Rocq has a separate `SProp` universe of proof-irrelevant propositions, and Agda has a predicative hierarchy `Propi` of such universes [Gilbert et al., 2019]. To use strict `Prop`, Rocq requires the flag `Allow StrictProp`, while Agda requires the option `{-# OPTIONS --prop #-}`.

Definitional proof irrelevance is useful to avoid having to prove equalities explicitly. Consider a relation on two naturals asserting the usual less-than relation, along with a type of bounded naturals whose boundedness can be forgotten.

$$\begin{aligned} \cdot \leq \cdot &: \text{Nat} \rightarrow \text{Nat} \rightarrow \text{Prop} \\ \text{BNat} &: \text{Nat} \rightarrow \text{Type}_0 \\ \text{bNat} &: \Pi(n \, m : \text{Nat}). n \leq m \rightarrow \text{BNat} \, m \\ \text{getNat} &: \Pi(m : \text{Nat}). \text{BNat} \, m \rightarrow \text{Nat} \\ \text{getNat} \, m \, (\text{bNat} \, n \, m \, p) &\rightsquigarrow n \end{aligned}$$

A desirable property of bounded naturals is that two bounded naturals are equal if their contained naturals are equal.

$$eqBNat : \forall(m : \text{Nat}) (b_1 \, b_2 : \text{BNat} \, m). \text{getNat} \, m \, b_1 \equiv \text{getNat} \, m \, b_2 \rightarrow b_1 \equiv b_2$$

If we try to prove this by destructing b_1 and b_2 as $(\text{bnat } n_1 \text{ } m \text{ } p_1)$ and $(\text{bnat } n_2 \text{ } m \text{ } p_2)$, $\text{getNat } m \text{ } b_1$ and $\text{getNat } m \text{ } b_2$ reduce to n_1 and n_2 .

$$\text{eqBNat } m \text{ } (\text{bnat } n_1 \text{ } m \text{ } p_1) \text{ } (\text{bnat } n_2 \text{ } m \text{ } p_2) : n_1 \equiv n_2 \rightarrow \text{bnat } n_1 \text{ } m \text{ } p_1 \equiv \text{bnat } n_2 \text{ } m \text{ } p_2$$

While we have an equality $n_1 \equiv n_2$, we do *not* have a proof of $p_1 \equiv p_2$ ¹. Depending on how the inequality predicate $\cdot \leq \cdot$ is implemented, it may be possible to prove propositionally that any two inequality proofs are equal. Alternatively, if propositions are definitionally proof irrelevant, the inequality proofs can be ignored. Then rewriting the goal by the given equality is sufficient for it to be proven by reflexivity.

$$\begin{aligned} \text{eqBNat } m \text{ } (\text{bnat } n_1 \text{ } m \text{ } p_1) \text{ } (\text{bnat } n_2 \text{ } m \text{ } p_2) \text{ } e & : \text{bnat } n_1 \text{ } m \text{ } p_1 \equiv \text{bnat } n_2 \text{ } m \text{ } p_2 \\ & := \text{rewrite } e \text{ in refl} \end{aligned}$$

Uniqueness of identity proofs (UIP). The *uniqueness of identity proofs* (UIP) asserts that inhabitants of the same propositional equality are themselves propositionally equal. It can be proven using *Axiom K* [Streicher, 1993], a computational eliminator for propositional equalities of type $a \equiv a$.

$$\begin{aligned} K : \forall (A : \text{Type}) (a : A) (P : a \equiv a \rightarrow \text{Prop}) (p : a \equiv a). P \text{ refl} \rightarrow P \text{ } p \\ K \text{ } A \text{ } a \text{ } P \text{ refl } d \rightsquigarrow d \end{aligned}$$

Agda’s default pattern matching behaviour, which permits matching on an equality of $a \equiv a$ as reflexivity, admits a proof of UIP as well as defining Axiom K. While not inherently part of Rocq’s type theory, Axiom K is axiomatized in the standard library as `Logic.Eqdep.eq_rect_eq`.

UIP is similarly useful to avoid reasoning about equalities between equalities when the only canonical proof of an equality is reflexivity, especially in settings without proof irrelevance. While UIP augments the reasoning power of the type theory, there are types whose equality proofs are already propositionally equal. In particular, if a type has decidable equality, *i.e.* $(x \equiv y) \vee \neg(x \equiv y)$ for any given x, y of that type, then its equalities are themselves equal [Hedberg, 1998].

Strong elimination. Destructing or eliminating an element of an inductive datatype into a type in a larger universe is known as *strong* or *large* elimination. That is, a term whose type is in Type_i is eliminated to return a term whose type is in Type_j for some $j > i$. For inductive proofs in Prop , this includes eliminating into any non-proposition type.

Strong elimination is a necessary ingredient in discriminating constructors of proof-relevant datatypes, such as the booleans. While `true` and `false` are syntactically distinct, proving their propositional inequivalence requires lifting the booleans to the propositions of truthhood $\top$ and falsehood $\perp$.

$$\begin{aligned} \text{lift} : \text{Bool} \rightarrow \text{Prop} \\ \text{lift} := \lambda b. \text{if } b \text{ then } \top \text{ else } \perp \end{aligned}$$

This branching expression is a strong elimination because it returns a type, or equivalently because its return type is Prop , a universe. Letting `cong f` be a proof of congruence of f

¹Technically, this should be a proof of the equality where p_1 has been rewritten by the equality $n_1 \equiv n_2$ so that it has the same type as p_2 .

over an equality, to complete the proof of $\text{true} \equiv \text{false} \rightarrow \perp$, given $e : \text{true} \equiv \text{false}$, the trivial proof of truthhood tt is rewritten by the lifted equality $\text{cong lift } e : \top \equiv \perp$.

```

trueNotFalse : true ≡ false → ⊥
trueNotFalse := λe. rewrite (cong lift e) in tt

```

2.2 Axioms

Rocq, Lean, and Agda all have mechanisms for defining axioms or postulates, which are declarations of constants without definitions. Although not all axioms are consistent, there are many well-studied axioms commonly used in practice that are worth considering as extensions in their own right. The axioms described here are included in many standard libraries, and proof developments can choose whether to import them.

Extensionality principles. Some models of type theory semantically equate things that are not syntactically (either definitionally or propositionally) equal; extensionality principles adds semantic equalities as propositional equalities. Examples include *function extensionality*, which equates two functions if they are pointwise equal, and *propositional extensionality*, which equates two propositions if they are biimplicated.

```

funext : ∀(A : Type) (B : A → Type) (f g : Π(x : A). B x). (∀(x : A). f x ≡ g x) → f ≡ g
propext : ∀(A B : Prop). (A ↔ B) → A ≡ B

```

These axioms are found in the standard libraries of Lean as `funext` and `propext`, and of Rocq as `Logic.FunctionalExtensionality.functional_extensionality` and `Logic.PropExtensionality.propositional_extensionality`. A notable consequence of propositional extensionality is propositional proof irrelevance.

Another example is *univalence*, which asserts an equivalence between propositional equality and equivalence, *i.e.* given two types A, B , the equivalence $(A \equiv B) \simeq (A \simeq B)$ holds. This is the core principle underlying Homotopy Type Theory (HoTT) [Univalent Foundations Program, 2013]. There are several ways to define equivalence; the idea is that it captures a propositionally proof-irrelevant isomorphism. Univalence together with proof irrelevance implies propositional extensionality, since biimplicated propositions are isomorphic by irrelevance, and univalence gives an equality from the isomorphism. Univalence alone also implies function extensionality by a more complex argument [Univalent Foundations Program, 2013, Chapter 4.9].

One application of function and propositional extensionality is encoding a function as a relation whose functionality is proven *a posteriori*. This is a frequent pattern in proof assistants, as inductive relations often have better ergonomic support than dependently-typed reasoning over functions. For example, consider a two-place predicate over a representation of types Ty and terms Tm , which has the type $Ty \rightarrow Tm \rightarrow \text{Prop}$. If we have trouble defining this predicate recursively due to termination issues or inductively due to strict positivity issues, we can instead view it as a function from Ty to a predicate $Tm \rightarrow \text{Prop}$ and try encoding it as a relation:

$$R : Ty \rightarrow (Tm \rightarrow \text{Prop}) \rightarrow \text{Prop}.$$

Such a relation could be a *logical relation* [Tait, 1967] used to model typed lambda calculi, where a Ty is interpreted as a set of Tms . Functionality of R demonstrates that Tys have unique interpretations. To show that R is functional, *i.e.*

$$\forall(A : Ty)(P Q : Tm \rightarrow \text{Prop}). R A P \rightarrow R A Q \rightarrow P \equiv Q,$$

it suffices to show that $\forall(a : Tm). P\ a \leftrightarrow Q\ a$, since $\forall(a : Tm). P\ a \equiv Q\ a$ follows from propositional extensionality, and finally $P \equiv Q$ from function extensionality.

The disadvantage of axiomatic equalities is that rewriting by them does not reduce, which can make reasoning about terms rewritten by such equalities difficult. There are type theories beyond MLTT and CIC that are designed so that these principles are instead provable theorems, such as cubical type theories [Bezem et al., 2019; Cohen et al., 2018; Angiuli et al., 2017, 2021] and Cubical Agda [Vezzosi et al., 2019] for univalence, and observational type theory [Altenkirch and McBride, 2006; Altenkirch et al., 2007; Pujet and Tabareau, 2022, 2023, 2024] for function and propositional extensionality.

Classical principles. There are a number of classical axioms that do not hold intuitionistically. The most common is the principle of *excluded middle* (EM), which asserts that all propositions are either true (inhabited) or false (uninhabited). EM is equivalent to several other principles, including *double negation elimination* (DNE), $\forall(A : \text{Prop}). \neg\neg A \rightarrow A$, and *Peirce’s law*, $\forall(A\ B : \text{Prop}). ((A \rightarrow B) \rightarrow A) \rightarrow A$. More powerful axioms which imply EM include the axiom of choice and the (in)definite description operators, which deal with extracting a concrete piece of data out of merely knowing that such a piece of data exists.

Because a large majority of mathematics is done classically, many communities freely use classical axioms when mechanizing mathematics. The axiom of excluded middle, for instance, is declared in Rocq as `Logic.Classical_Prop.classic`, and in Lean as `em`. The `Logic` subdirectory of Rocq’s standard library contains classical axioms along with proofs about their properties. Similarly, Lean’s mathematical library `mathlib` [mathlib Community, 2020] contains proofs that rely on classical axioms, and tactics such as `tauto` automatically apply classical reasoning.

2.3 Extensions and inconsistencies

One has to be careful not to choose a set of features and axioms that render the type theory logically inconsistent and thus useless for proving. A number of them are known to be incompatible with one another; below are a few such combinations.

- Strong elimination is inconsistent for large impredicative inductives. Hook and Howe [1986] show that impredicative dependent pairs with pair projections, which correspond to strong elimination, are inconsistent. Coquand [1992] also demonstrates the inconsistency of strongly eliminating a large inductive type U with a single constructor of type $\forall(X : \text{Prop}). (X \rightarrow U) \rightarrow U$.
- Strong elimination is also inconsistent for inductive propositions when `Prop` is proof irrelevant. As seen above, strong elimination suffices to show that $\neg(\text{true} \equiv \text{false})$. If `Bool` is defined in a proof-irrelevant `Prop`, $\text{true} \equiv \text{false}$ would hold by definition, which is a contradiction.
- Strong elimination is once again inconsistent for inductive propositions in the presence of impredicativity and classical principles such as excluded middle. `Logic.Berardi` in Rocq’s standard library, which is a modern implementation of the construction by Barbanera and Berardi [1996], uses excluded middle to derive propositional proof irrelevance, which can be used as above to derive a contradiction.
- UIP is inconsistent with univalence. Intuitively, univalence produces an equality between two types given an equivalence between them, and there are types that are equivalent in multiple, provably different ways, so there are equalities between them that are provably different, thus violating UIP. Concretely, `Bool` is equivalent to itself

in two different ways, either by mapping booleans to themselves or to their negation, so there are distinct proofs of $\text{Bool} \equiv \text{Bool}$ [Univalent Foundations Program, 2013, Example 3.1.9].

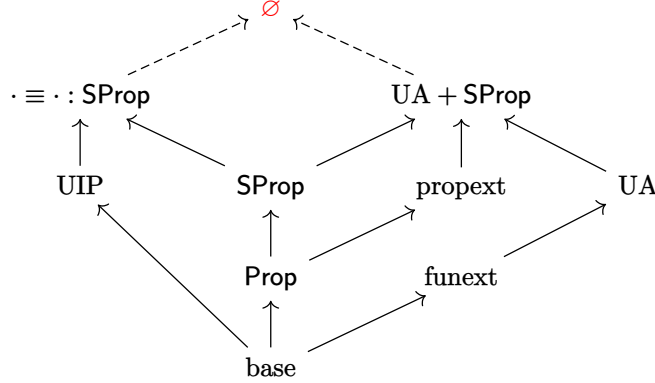


Figure 1: A compatibility graph of theories with impredicative **Prop**, proof irrelevance (**SProp**), UIP, univalence (**UA**), function extensionality (**funext**), propositional extensionality (**propext**), and (in)compatible combinations.

Figure 1 illustrates some of these relationships between extensions. The arrows point from one theory to a strictly more expressive one; for instance, a theory with propositional extensionality extends the equalities of a theory containing a universe of propositions, and a theory with univalence can derive function extensionality. At the top of the graph, the dotted arrows indicate the incompatibility of a theory that implies UIP with one that contains univalence: there is no possible encompassing theory.

To prevent inconsistencies, proof assistants hide features behind option flags or disallow them entirely. Rocq, Lean, and Agda generally disallow strong elimination for inductive propositions in **Prop** and **SProp**. The exceptions are *syntactic subsingletons*, which are inductives that syntactically have at most one inhabitant, such as $\top$, $\perp$, and conjunction of propositions. While Rocq’s impredicative **Prop** universe is not proof irrelevant, Rocq still forbids strong elimination even for inductives that are not large to allow the use of classical principles. Its compiler flag `-impredicative-set` enables impredicativity for its **Set** universe while still allowing strong elimination of small impredicative inductives, but this flag is not well supported.² In Agda, the `--with-K` flag enables Axiom K and the `--without-K` flag disables it, while the features enabled by Cubical Agda using the `--cubical` flag prove univalence as a theorem. There is also a `--safe` flag which disallows, among other combinations, having both `--with-K` and `--cubical`.

Proof assistants provide some limited ability in tracking the usage of features and axioms. In Rocq, the axioms and unsafe flags used by a definition can be listed using the command `Print Assumptions`, while in Lean, the axioms can be listed using `#print axioms`. In Agda, along with checking for inconsistent option combinations, `--safe` ensures the absence of any `postulates`. Because option flags can be enabled at module-level granularity, Agda also has notion of *(co)infective* flags: an infective flag used in one module must be used by all modules that depend on that module, while a coinfective flag used in one module may only depend on modules that also use that flag.

While users of these proof assistants have been getting along fine with these compiler

²<https://github.com/coq/coq/issues/9458>

tools for tracking features and axioms for the past decade or more, there is plenty of room of improvement; I have identified three shortcomings these systems.

- There is no way of asserting against an extension; that is, there is no general mechanism to tell the type checker to fail if a particular definition uses some feature or axiom. Such an assertion would guarantee that definition safe to be used by others that use an incompatible extension. For instance, ensuring that a proof does not use Axiom K means that it may be used by another proof that does use univalence. In this specific case, Agda does have the `--without-K` flag for this purpose, but few options have a corresponding anti-option. Meanwhile, Rocq and Lean’s axiom printing mechanism does not modify type checking behaviour.
- The scope of feature flags is too coarse. They range from project-level compiler flags down to module-level option flags, but even this level of granularity is not the appropriate one: modules are intended for organizing definitions by semantic content, rather than by the collection of features they happen to all use. This prevents reuse of a definition in one module inside another module if they happen to have incompatible features, even if that particular definition does not depend on any features at all.
- Completely disallowing an incompatible extension is unnecessarily restrictive. For example, when both `--safe` and `--with-K` are enabled in Agda, definitions from modules with `--cubical` cannot be mentioned anywhere. However, it should be acceptable to use a theory with UIP to state and prove properties of univalence, while never eliminating an equality derived from univalence. To generalize, definitions from an incompatible extension should be permitted in the type of a term but restricted in the term itself, so that we can talk *about* an extension without *using* it.

An ideal system for enabling and disabling features and axioms, then, should track which ones are and aren’t used, at least at the definition level, and distinguish between mentioning and using them. This suggests that a system for dependency tracking would be a good starting point. More precisely, the Dependent Calculus of Indistinguishability (DCOI) [Liu et al., 2024b], a type system that incorporates dependency tracking and dependent types, could be a suitable framework for tracking the usage of extensions within type theories.

3 A primer on DCOI

DCOI is a Pure Type System (PTS) [Barendregt, 1991] augmented with dependency tracking [Abadi et al. [1999]. DCOI^ω [Liu et al., 2025] is a logically consistent instantiation of DCOI’s PTS rules and axioms with a predicative universe hierarchy, making it suitable as a foundation for theorem proving. In a typing judgement, dependency tracking appears as annotations both on the variables in the context, to indicate how they may be used by the term being typed, and alongside the type of the term, to indicate how the term itself may be used. As an example, consider the following typing judgement for a constant function.

$$A :^H \text{Type} \vdash \lambda x^L y^H. x :^L A^L \rightarrow A^H \rightarrow A$$

The concrete levels used here are low (L) and high (H) where $L < H$. In the context of information flow, these levels correspond to low- and high-security computations where low-security computations may not inspect the values of high-security ones. They can also be thought of in terms of computational irrelevance, where something marked as computationally irrelevant (H) must not play a part in the execution of relevant programs (L), and may even be erased away after compilation.

This constant function at low, which returns its first argument x and ignores its second argument y , must therefore mark x as low to return it. Marking y as high guarantees that the function may not return it. While the body of the low function cannot return a high argument, its *type* can depend on a high term, demonstrated by the high-annotated type A in the context, which is used in the type of the function. The intuition is that A does not play a part in the run-time execution of the constant function, but is otherwise permitted to participate in compile-time type checking.

Other terms we may wish to mark as high are proofs to guarantee that a compiler can erase them away after type checking. For instance, a length function on a computationally relevant list would return a relevant natural; in contrast, a proof that no relevant natural is strictly less than zero would be marked as irrelevant.

$$\begin{aligned} \text{length} &: ^L \Pi(A : ^H \text{Type}) (l : ^L \text{List } A). \text{Nat} \\ \text{zeroNotGt} &: ^H \forall(n : ^L \text{Nat}) (p : ^H n < 0). \perp \end{aligned}$$

Though these examples only use two levels L and H , any instance of DCOI is parametrized over some meet-semilattice, from which its dependency levels are drawn. Rather than tracking relevance or information flow, the goal is to track extensions using a lattice that associates dependency levels to sets of added extensions. Before discussing this mapping in [Section 4](#), I summarize briefly the key aspects and properties of DCOI that highlight how dependency tracking interacts with terms and typing.

Relative relevance. The intuition of computational relevance and irrelevance is not fixed to the low and high levels, but is a relative concept between any two ordered dependency levels. Suppose there is a super-high level S such that $L < H < S$. Then just as a low term may not meaningfully use a high term, a high term also may not meaningfully use a super-high term. The following typing judgement demonstrates an application of these three levels.

$$P : ^H \text{Nat}^S \rightarrow \text{Prop}, n : ^S \text{Nat} \vdash \lambda p.^L p : ^L (P \ n^S)^L \rightarrow P \ (n + 1)^S$$

In the context, n is a super-high natural, and P is a high predicate which takes as argument a super-high natural. Once again, the term being typed is a low function, while higher terms are involved in its type. Although the function is an identity function, its domain and codomain types are syntactically different applications of P . This judgement holds due to *indistinguishability*.

Indistinguishability. In general, if $\ell_0 < \ell_1$, then at observer level ℓ_0 , the function application $f \ x^{\ell_1}$ must be equal to $f \ y^{\ell_1}$ regardless of what x and y are. We say that they are *indistinguishable* at level ℓ_0 , written as $\cdot =^{\ell_0} \cdot$. Below are the two different rules for indistinguishability of a function application depending on the observer level ℓ_0 and the argument level ℓ_1 .

$$\begin{array}{c} \text{I-APP-INDIST} \\ \frac{\Gamma \vdash f =^{\ell_0} g \quad \ell_0 < \ell_1}{\Gamma \vdash f \ a^{\ell_1} =^{\ell_0} g \ b^{\ell_1}} \end{array} \qquad \begin{array}{c} \text{I-APP-DIST} \\ \frac{\Gamma \vdash f =^{\ell_0} g \quad \Gamma \vdash a =^{\ell_1} b \quad \ell_0 \not< \ell_1}{\Gamma \vdash f \ a^{\ell_1} =^{\ell_0} g \ b^{\ell_1}} \end{array}$$

Indistinguishability plays the rôle of definitional equality when checking whether two types are the same. This is why the above example type checks, since $P \ n^S =^H P \ (n + 1)^S$ holds.

Similarly, letting k be the the low-level constant function above, $k x^L y^H =^L k x^L z^H$ holds, which expresses the idea that k is truly constant in its second argument.

DCOI internalizes indistinguishability by indexing its propositional equality type with an observer level, reflected by the rule for reflexivity below. In particular, the propositional equality $k x^L y^H \equiv^L k x^L z^H$ is provable by reflexivity since the two sides are already indistinguishable at low, the observer level of the equality.

$$\frac{\Gamma \vdash a =^{\ell_0} b}{\Gamma \vdash \text{refl} :^{\ell} a \equiv^{\ell_0} b}$$

Elimination of higher falsehoods. The principle that lower-level terms may not meaningfully depend on higher-level terms means that destructors that produce lower-level terms may not destruct higher-level terms. This holds even if the term being destructed contains no inner information (such as $\top$ or an equality proof), since reducing the destruction on a constructor requires knowing whether the term being destructed is a constructor at all.

The sole exception is the eliminator for $\perp$, since it has no constructors, so there is no information to reveal. It is well typed at any level independent of its target b or its type A .

$$\frac{\Gamma \vdash b :^{\ell_0} \perp \quad \Gamma \vdash A :^{\ell_1} \text{Type}_i}{\Gamma \vdash \text{absurd } b :^{\ell} A}$$

The computational interpretation of having a proof of $\perp$ to eliminate is that we have reached an impossible dead branch, so what we do with it never matters since it never executes. For example, we can define a function that extracts the head of a nonempty list by using `absurd` to handle the impossible empty list case.

$$\begin{aligned} \text{head} &: ^L \Pi(A : ^H \text{Type})(l : ^L \text{List } A) (p : ^H 0 < \text{length } l^L). A \\ \text{head } A \text{ nil } p &:= \text{absurd } (\text{zeroNotGt } 0^L) \\ \text{head } A (\text{cons } x \text{ xs}) p &:= x \end{aligned}$$

Subsumption and downgrading. While lower-level terms cannot inspect higher-level terms, higher-level terms can inspect lower-level terms. A lower-level term can also be raised to a higher level by *subsumption*: if a term is well typed at level ℓ_0 , then it is also well typed with the same type at a higher level $\ell_1 \geq \ell_0$. This admissible rule is given below on the left.

$$\begin{array}{c} \text{SUBSUMPTION} \\ \frac{\Gamma \vdash a :^{\ell_0} A \quad \ell_0 \leq \ell_1}{\Gamma \vdash a :^{\ell_1} A} \end{array} \qquad \begin{array}{c} \text{DOWNGRADE} \\ \frac{\Gamma \vdash a \equiv^{\ell_1} b \quad \ell_0 \leq \ell_1}{\Gamma \vdash a \equiv^{\ell_0} b} \end{array}$$

However, if two terms are indistinguishable by some observer level ℓ_1 , then they are indistinguishable by a *lower* observer level ℓ_0 by *downgrading*, given above on the right. From a security flow perspective, the higher the observer level, the more secure values may be observed, and the more things are distinguishable, since more secure values need to be compared instead of being ignored. Going down an observer level means more things are being hidden away, so more values appear to be indistinguishable from one another.

4 Lattices of type theories

The key premise of using DCOI for extension tracking is associating dependency levels with various type theories, creating a lattice whose points are the sets of extensions enabled in

each theory, ordered by subset. We begin with a bottom dependency level for the base type theory corresponding to the empty set of extensions. For each additional construct corresponding to a feature or axiom, we add a new dependency level above the theory it extends. For instance, there could be a K eliminator that type checks at a level for UIP, or a built-in excluded middle axiom that type checks at a level for classical reasoning.

$$\text{K } P \text{ } p \text{ } d :^{\text{uip}} P \text{ } p \qquad \text{em } A :^{\text{cl}} A \vee \neg A$$

Because level annotations are part of contexts and typing judgements, well-typedness of a particular definition also specifies exactly where it is safe to be used, guaranteeing that it never exploits an extension without permission. A definition that can be typed at the bottom level would be safe to be used at all levels by subsumption, and guaranteed to never employ, say, classical reasoning. Indistinguishability reflects this guarantee, as it asserts the property that uses of terms from higher forbidden theories can only be trivial, such as ignoring them or passing them around uninspected.

As dependency levels form a meet-semilattice, any two theories must have a meet (*i.e.* intersection), which corresponds to only the constructs that they both have in common, and which are therefore safe to use in either theory. If the join (*i.e.* union) of two theories exist, then the constructs introduced in either one can be used at the joined level. Crucially, not all joins exist; a UIP level cannot be joined with a univalence level, since their coexistence is contradictory. The shape of the lattice depends on the compatibilities between theories, as well as the implication order of extensions, since one theory that encompasses the consequences of another can be placed above that other theory. The compatibility graph in [Figure 1](#) is an example of a concrete lattice of theories, where the arrows point towards the greater theory and indicate the direction in which definitions can be raised.

The property that the type of a term can itself be well typed at a level independent of the term's permits proving propositions about an incompatible theory without causing an inconsistency. For example, we can assert the computational behaviour of the K eliminator even in a theory with univalence.

$$\text{refl} :^{\text{ua}} \text{K } P \text{ refl } d \equiv^{\text{uip}} d$$

Following the rules of DCOI, each individual theory must be logically consistent. If an inconsistency exists at any theory, by the elimination of higher falsehoods, the inconsistency propagates to all lower theories, including the bottom theory. Then by subsumption, the inconsistency at the bottom theory can be raised to propagate to all higher theories, and the entire lattice would be inconsistent. This means that if eliminating falsehoods works exactly as in DCOI, any theory that features nontermination would not be permitted.

As the goal is to exclude incompatible extensions from a proof assistant, disallowing logically inconsistent theories is a desirable trait. Nevertheless, there may be a few ways to modify falsehood elimination to permit them. One way is to disallow eliminating falsehoods to lower levels, allowing it only to the same level. Another is to take ideas from works from the Trellys project, such as λ^θ [Casinghino et al., 2014] and Sep³ [Kimmell et al., 2012], and impose a value or termination restriction on falsehoods being eliminated. If the falsehood in an inconsistent theory is nonterminating or not a value, then it cannot be eliminated at all, preventing its propagating to lower theories.

One catch is that a theory whose extension is a new definitional equality (*i.e.* a new rule for indistinguishability) cannot be contained within its level. Even if that equality is defined for a particular observer level, it will hold for all lower observer levels by downgrading, and the extension will be available to all lower theories. This effect cannot be mitigated

using restrictive premises, as violating downgrading violates many other desirable properties, including transitivity of definitional equality [Liu et al., 2025].

Consequently, adding strictness to an existing **Prop** universe rather than adding an entirely separate **SProp** universe is not possible, as the type checker can lift two proofs at a non-proof-irrelevant level up to the proof-irrelevant level and equate them there. One solution is to invert the conventional order and place non-strict **Prop** *above* strict **Prop**, so that *disabling* strictness is an extension. The inversion slightly complicates the lattice in Figure 1, as the level with propositional UIP via the K eliminator would be above the level with propositional equality in non-strict **Prop**, which in turn would be above the level with equality in strict **Prop** and thus with definitional UIP. Such a lattice has the unusual property that the univalent level would be joinable with the non-strict **Prop** level, but not with the UIP or strict **Prop** levels above and below it.

4.1 Objectives

This project should answer the following questions:

1. What kinds of extensions would fit within this framework? Some broad classifications of extensions might be ones that add new type universes (*e.g.* **SProp**), ones that expand the rules for existing constructs (*e.g.* impredicativity, strong elimination), ones that add new computational constructs with reduction rules (*e.g.* Axiom K), and ones that add new axiomatic constructs without reduction rules (*e.g.* function and propositional extensionality, excluded middle).
2. What are useful applications of being able to freely refer to other theories? The previous example that referred to the K eliminator, while true, is stating a trivial fact that is already provable within the UIP theory. Are there meaningful theorems about one theory that cannot be proven in that theory, but can be proven in a different yet potentially incompatible theory?
3. How would a particular lattice of theories be modelled to show logical consistency? Ideally, the technique used to model a particular lattice should be broadly applicable and sufficiently extensible to a different lattice without redoing all the work, so that adding more extensions remains sustainable.

To answer these questions, the project would be divided into two portions. The first is an implementation of a type checker for a specific lattice of type theories. The lattice should contain a sufficiently diverse set of labels and their orders. Figure 1 is a good place to start, as it contains theories in different classifications with different interactions.

To evaluate the viability of such a type checker, a standard library would be implemented to exercise all levels of the lattice. The standard libraries of Rocq³, Agda⁴, and Lean⁵ are good sources for inspiration, as many of their files use the features and axioms mentioned in Section 2. An implementation would also serve to verify which extensions are indeed invalid by demonstrating the inconsistencies they yield.

Additionally, writing and checking large proofs would be feasible, which helps with exploring more complex applications. One example is computing *semisimplicial types*, which

³<https://coq.inria.fr/distrib/current/stdlib/>

⁴<https://agda.github.io/agda-stdlib/master/>

⁵https://leanprover-community.github.io/mathlib4_docs/

is an open problem in HoTT. Two-level type theory (2LTT) [Altenkirch et al., 2016; Anenkov et al., 2023] is one solution that contains separate theories with univalence and with UIP, and Agda implements 2LTT as the `--two-level` extension. We can use libraries⁶ for 2LTT in Agda as inspiration for testing whether this project’s implementation of orthogonal univalence and UIP extensions can handle the same proof load. Section 6.1 discusses some more details about 2LTT.

The useability of the implementation would inform the design of the system, such as level inference and level polymorphism. Annotating definitions and arguments with every single extension it uses is an unreasonable burden on a proof assistant user. The prototype implementation accompanying DCOI [Liu et al., 2024a] therefore has rudimentary level inference, defaulting to a minimum level. However, the prototype’s lattice is the usual order of the naturals, and its examples typically use no more than two levels. An implementation with a lattice containing incompatible extensions and applications that involve more than two theories would pinpoint what is required from level inference in practice for a more sophisticated inference algorithm.

Rewriting libraries in this implementation is also an exercise in determining where code duplication occurs and whether level polymorphism would help eliminate it. Although subsumption allows lifting definitions vertically, so to speak, from lower theories to higher ones, it does not allow transporting definitions horizontally from one theory to a different, incompatible theory. In addition, while a function can quantify over terms at specific levels, it cannot quantify over *all* levels, nor over levels that satisfy some ordering constraint. Having a library of examples in the implementation would reveal whether these are real concerns to be addressed and what kind of level polymorphism could address them.

The second portion is a formalized and ideally mechanized proof of consistency. Because consistency is a semantic property and depends on the strength of the metatheory used to model the type theory, the formalization should model a lattice with (at least at first) only one level above the base theory, the simplest nontrivial lattice. The focus would be on how to combine two different models of type theory, not on accommodating as many as possible from the outset.

A sensible starting point would be taking the mechanization of DCOI^ω [Liu et al., 2025], which proves consistency and normalization of what would be the base theory in the lattice, and picking a reasonable feature to extend it with. UIP and function extensionality are good candidates, as they are expected to already hold in the semantic model. In this case, the challenge is to prove that the base theory does *not* prove UIP or function extensionality to demonstrate that other incompatible extensions could be added to the base theory.

Moreover, the semantic model is a syntactic logical relation indexed by well-founded universe levels, which may limit its extensibility; it cannot be straightforwardly extended to accommodate impredicativity, nor to accommodate typed definitional equality. If the PTS rules and axioms of DCOI are instead instantiated to a single impredicative universe with no universe hierarchy or strong elimination, we can adapt the logical relation for the Calculus of Constructions by Geuvers [1995] to prove consistency and normalization.

A possible alternative is to use *syntactic* modelling [Boulier et al., 2017], which would involve a type-preserving translation into another type theory whose consistency is well established, guaranteeing consistency of the original system. While there exist syntactic models of other type theories [Gilbert et al., 2019; Winterhalter, 2024] with notions of irrelevance, which is one application of indistinguishability, a syntactic model of dependency tracking with dependent types is unexplored.

⁶e.g. <https://github.com/UnivalencePrinciple/2LTT-Agda>

5 Prior work

This project builds on prior work on DCOI [Liu et al., 2024b] and DCOI^ω [Liu et al., 2025], on both of which I am second author. For the former paper, I implemented a prototype type checker for DCOI augmented with inductive types by extending the minimal dependent type checker `pi-forall` [Weirich, 2022], and I wrote examples using the type checker and motivating examples for DCOI. I also proved a few of the lemmas in the mechanization. For the latter paper, I wrote about half of the prose, mostly for the earlier sections, and proved a few of the lemmas as well. As part of an investigation toward incorporating a relational model for DCOI, I mechanized a PER model for MLTT based on the logical relation used to prove consistency of DCOI^ω , but ultimately the gap between MLTT and DCOI could not be bridged, so this work does not appear in the final paper.

Outside of DCOI, I have worked on Stratified Type Theory (StratTT) [Chan and Weirich, 2025], which annotates typing judgements similarly to dependency tracking, but the annotations are universe levels. Instead of stratifying universes into a hierarchy, typing judgements themselves are stratified, and there is a single universe whose type is itself. To ensure consistency in the presence of this *type-in-type* rule, dependent functions may only quantify over types at strictly lower levels, which enforces predicativity. Although StratTT is not a dependency tracking system in the same way DCOI is, it demonstrates that there may be multiple ways to retain usage information that enforces consistency. Even if the particular setup for DCOI turns out not to be suitable for this project, it may be reasonable to instead explore a more StratTT-like structure.

6 Related work

6.1 Multi-system frameworks

Two-level type theory. The most similar work to extension tracking is two-level type theory (2LTT) [Altenkirch et al., 2016; Annenkov et al., 2023]. It consists of an inner homotopical type theory with univalence and an outer intensional type theory with UIP, along with a conversion operation $\uparrow \cdot$ from the inner theory to the outer. The inner and outer type theories have independent type formers, including separate inner (path) equality types $\cdot \equiv^p \cdot$ and outer (strict) equality types $\cdot \equiv^s \cdot$. Converting an inner equality does *not* yield the outer equality type; otherwise, univalence on inner equalities could be converted to univalence on outer equalities, which would contradict UIP of the outer equality.

These inner and outer levels are different from dependency levels in DCOI, where all levels share the same type formers, and a lower equality can be raised to a higher equality by subsumption. In particular, if a lattice of type theories includes one that supports UIP, then that level will prove that proofs of the same equality at *all* levels are equal. Meanwhile, in 2LTT, UIP only holds for proofs the outer equality and not for converted proofs of the inner equality.

The conversion operator can be thought of as an explicit subsumption, and how it interacts with the inner and outer equalities is similar to how DCOI’s propositional equality interacts with indistinguishability at lower and higher levels. To demonstrate, given two inner terms x, y , the implication $\uparrow x \equiv^s \uparrow y \rightarrow x \equiv^p y$ holds in 2LTT while the converse generally does not. Similarly, in DCOI, $x \equiv^H y \rightarrow x \equiv^L y$ holds by downgrading while the converse also generally does not.

The Trellys project. Instead of combining multiple type theories, the Trellys project focussed on combining dependently-typed logical reasoning with (potentially nonterminating) functional programming. The main works within the project are λ^θ [Casinghino et al., 2014], which classifies typing judgements of a single language into logical and programmatic fragments; Zombie [Sjöberg, 2015], an implementation that closely follows λ^θ ; Sep³ [Kimmell et al., 2012], which syntactically separates proofs from programs; and Nax [Ahn, 2014], which augments dependent types with Mendler-style recursion schemes.

Of these three, DCOI is closest to λ^θ , whose logical and programmatic classifications are similar to DCOI’s dependency levels. Because the logical fragment is subsumed within the programmatic fragment, the additional features found in the programmatic fragment can be thought of as an extension of the logical one. Notable features of the extension include isorecursive types and unrestricted recursion, allowing for nonterminating programs. Normalization of the logical fragment is ensured by only allowing boxed programs to be applied to its functions, only allowing unboxing of values, and restricting reduction to call by value. While the boxing mechanism is similar to level annotations on function domains in DCOI, the value restrictions are specific to handling the presence of potential divergence.

Casinghino [2015] extends the work done for λ^θ by looking at a number of languages with logical and programmatic fragments, beginning with a simply typed calculus with classifications (also named λ^θ), extending it with types dependent on terms in LF^θ , then further extending that with terms dependent on types (*i.e.* type polymorphism) and types dependent on types (*i.e.* type-level computation) in PCC^θ . The proofs of normalization for these calculi use partially step-indexed logical relations, where stepping only occurs in the programmatic fragment. However, the proof technique does not scale up to handle large elimination or a universe hierarchy, both of which are present in Zombie.

System DE. Similarly to PCC^θ , System DE [Liu and Weirich, 2023] is a dependent type system with logical and programmatic fragments, where the logical fragment is again normalizing. It extends System DC [Weirich et al., 2017] and is designed to be suitable as a core calculus of the Glasgow Haskell Compiler [GHC Development Team, 2004]; as such, it deals with explicit coercion proofs. In contrast to PCC^θ , System DE does not permit using a programmatic value as a logical value, so its logical relation proof does not require partial step-indexing to interpret the programmatic fragment. Consequently, System DE is able to include the type-in-type rule and large elimination in the programmatic fragment.

6.2 Other proof assistants

Section 2 broadly covers a number of optional features and common axioms in Rocq, Lean, and Agda. There are many other proof assistants of varying relevance not discussed above.

Idris 2 [Brady, 2021] is a dependently typed programming language with partiality. Definitions can be marked as `total`, `covering`, or `partial`; setting aside type-in-type, totality ensures consistency, covering ensures type safety while allowing divergence, and partiality does not ensure either. Because partiality subsumes covering subsumes totality, these modifiers can also be thought of as members of a lattice. Similarly to PCC^θ , Idris 2 is also call by value, and the partiality modifier is designed so that diverging terms can appear in types and be reasoned about while not being reduced during type checking. Although Idris 2 is based on Quantitative Type Theory [Atkey, 2018], it has no complete formal description that describes all of its core features, especially as it is a rapidly evolving language.

F* [Swamy et al., 2016] is a proof assistant with dependency tracking for different effects. Its dependency levels include a `Tot` level for total programs and a `Dv` level above it for

potentially diverging programs. In contrast to PCC^θ , the total fragment of F^* may not refer to the diverging fragment, so the proof of weak normalization of the total fragment involves a logical relation that does not consider levels above Tot . All the effects in F^* are implemented as indexed monads, which get compiled away to the new core calculus TotalF^* [Rastogi et al., 2021]; divergence aside, effects do not extend the internal type system.

7 Conclusion

In this report, I have described a number of common extensions to proof assistants that alter their core type theory when enabled. These proof assistants have mechanisms for tracking extensions and ensuring that incompatible extensions may not be used together. However, extension tracking is external to the type system, and internalization opens up opportunities for greater precision and expressivity in specifying extension usage. Such an internalization of extension tracking is reminiscent of dependency analysis, and in particular of the Dependent Calculus of Indistinguishability (DCOI), which tracks usage of terms at different dependency levels. Although primary applications of DCOI are information flow and irrelevance, I use it as inspiration for a framework that tracks extensions in the same way, incorporating multiple type theories in one system. Much work lies ahead to discover the practical expressivity of such a system, as well as its logical consistency, for its viability as a core for a proof assistant with extensions.

References

- Martín Abadi, Anindya Banerjee, Nevin Heintze, and Jon G. Riecke. 1999. A Core Calculus of Dependency. In *Proceedings of the 26th ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages* (San Antonio, Texas, USA) (POPL '99). Association for Computing Machinery, New York, NY, USA, 147–160. doi:10.1145/292540.292555
- Ki Yung Ahn. 2014. *The Nax Language: Unifying Functional Programming and Logical Reasoning in a Language based on Mendler-style Recursion Schemes and Term-indexed Types*. Ph.D. Dissertation. Portland State University. doi:10.15760/etd.2086
- Thorsten Altenkirch, Paolo Capriotti, and Nicolai Kraus. 2016. Extending Homotopy Type Theory with Strict Equality. In *25th EACSL Annual Conference on Computer Science Logic (CSL 2016) (Leibniz International Proceedings in Informatics (LIPIcs), Vol. 62)*, Jean-Marc Talbot and Laurent Regnier (Eds.). Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl, Germany, 21:1–21:17. doi:10.4230/LIPIcs.CSL.2016.21
- Thorsten Altenkirch and Conor McBride. 2006. Towards Observational Type Theory. <http://strictlypositive.org/ott.pdf>
- Thorsten Altenkirch, Conor McBride, and Wouter Swierstra. 2007. Observational equality, now!. In *Proceedings of the 2007 Workshop on Programming Languages Meets Program Verification* (Freiburg, Germany) (PLPV '07). Association for Computing Machinery, New York, NY, USA, 57–68. doi:10.1145/1292597.1292608
- Carlo Angiuli, Guillaume Brunerie, Thierry Coquand, Robert Harper, Kuen-Bang Hou (Favonia), and Daniel R. Licata. 2021. Syntax and models of Cartesian cubical type theory. *Mathematical Structures in Computer Science* 31, 4 (2021), 424–468. doi:10.1017/S0960129521000347
- Carlo Angiuli, Kuen-Bang (Favonia) Hou, and Robert Harper. 2017. Computational Higher Type Theory III: Univalent Universes and Exact Equality. doi:10.48550/arXiv.1712.01800
- Danil Annenkov, Paolo Capriotti, Nicolai Kraus, and Christian Sattler. 2023. Two-level type theory and applications. *Mathematical Structures in Computer Science* 33, 8 (2023), 688–743. doi:10.1017/S0960129523000130
- Robert Atkey. 2018. Syntax and Semantics of Quantitative Type Theory. In *Proceedings of the 33rd Annual ACM/IEEE Symposium on Logic in Computer Science* (Oxford, United Kingdom) (LICS '18). Association for Computing Machinery, New York, NY, USA, 56–65. doi:10.1145/3209108.3209189
- Franco Barbanera and Stefano Berardi. 1996. Proof-irrelevance out of excluded-middle and choice in the calculus of constructions. *Journal of Functional Programming* 6, 3 (1996), 519–526. doi:10.1017/S0956796800001829
- Henk Barendregt. 1991. Introduction to generalized type systems. *Journal of Functional Programming* 1, 2 (1991), 462–490. doi:10.1017/s0956796800020025
- Marc Bezem, Thierry Coquand, and Simon Huber. 2019. The Univalence Axiom in Cubical Sets. *Journal of Automated Reasoning* 63 (2019), 159–171.

- Simon Boulier, Pierre-Marie Pédro, and Nicolas Tabareau. 2017. The next 700 syntactical models of type theory. In *Proceedings of the 6th ACM SIGPLAN Conference on Certified Programs and Proofs* (Paris, France) (*CPP 2017*). Association for Computing Machinery, New York, NY, USA, 182–194. doi:10.1145/3018610.3018620
- Edwin Brady. 2021. Idris 2: Quantitative Type Theory in Practice. In *35th European Conference on Object-Oriented Programming (ECOOP 2021) (Leibniz International Proceedings in Informatics (LIPIcs), Vol. 194)*, Anders Möller and Manu Sridharan (Eds.). Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl, Germany, 9:1–9:26. doi:10.4230/LIPIcs.ECOOP.2021.9
- Chris Casinghino. 2015. *Combining Proofs and Programs*. Ph.D. Dissertation. University of Pennsylvania. <https://repository.upenn.edu/handle/20.500.14332/28018>
- Chris Casinghino, Vilhelm Sjöberg, and Stephanie Weirich. 2014. Combining proofs and programs in a dependently typed language. In *Proceedings of the 41st ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages* (San Diego, California, USA) (*POPL '14*). Association for Computing Machinery, New York, NY, USA, 33–45. doi:10.1145/2535838.2535883
- Jonathan Chan and Stephanie Weirich. 2025. Stratified Type Theory. In *Programming Languages and Systems*, Viktor Vafeiadis (Ed.). Springer Nature Switzerland, Cham.
- Cyril Cohen, Thierry Coquand, Simon Huber, and Anders Mörtberg. 2018. Cubical Type Theory: A Constructive Interpretation of the Univalence Axiom. In *21st International Conference on Types for Proofs and Programs (TYPES 2015) (Leibniz International Proceedings in Informatics (LIPIcs), Vol. 69)*, Tarmo Uustalu (Ed.). Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl, Germany, 5:1–5:34. doi:10.4230/LIPIcs.TYPES.2015.5
- The Coq Development Team. 2022. The Coq Proof Assistant. doi:10.5281/zenodo.5846982
- Thierry Coquand. 1992. The paradox of trees. *BIT Numerical Mathematics* 32 (March 1992), 10–14. doi:10.1007/BF01995104
- Leonardo de Moura, Soonho Kong, Jeremy Avigad, Floris van Doorn, and Jakob von Raumer. 2015. The Lean Theorem Prover (System Description). In *International Conference on Automated Deduction (Lecture Notes in Computer Science, Vol. 9195)*. Springer, Cham, Cham, Switzerland, 378–388. doi:10.1007/978-3-319-21401-6_26
- Herman Geuvers. 1995. A short and flexible proof of strong normalization for the calculus of constructions. In *Types for Proofs and Programs*, Peter Dybjer, Bengt Nordström, and Jan Smith (Eds.). Springer Berlin Heidelberg, Berlin, Heidelberg, 14–38.
- GHC Development Team. 2004. *The Glasgow Haskell Compiler*. The University Court of the University of Glasgow. <https://www.haskell.org/ghc/>
- Gaëtan Gilbert, Jesper Cockx, Matthieu Sozeau, and Nicolas Tabareau. 2019. Definitional proof-irrelevance without K. *Proc. ACM Program. Lang.* 3, POPL, Article 3 (Jan. 2019), 28 pages. doi:10.1145/3290316
- Michael Hedberg. 1998. A coherence theorem for Martin-Löf’s type theory. *Journal of Functional Programming* 8, 4 (1998), 413–436. doi:10.1017/S0956796898003153

- James G. Hook and Douglas J. Howe. 1986. *Impredicative Strong Existential Equivalent to Type:Type*. Technical Report TR86-760. Cornell University. <https://hdl.handle.net/1813/6600>
- Garrin Kimmell, Aaron Stump, Harley D. Eades, Peng Fu, Tim Sheard, Stephanie Weirich, Chris Casinghino, Vilhelm Sjöberg, Nathan Collins, and Ki Yung Ahn. 2012. Equational reasoning about programs with general recursion and call-by-value semantics. In *Proceedings of the Sixth Workshop on Programming Languages Meets Program Verification* (Philadelphia, Pennsylvania, USA) (*PLPV '12*). Association for Computing Machinery, New York, NY, USA, 15–26. doi:10.1145/2103776.2103780
- Yiyun Liu, Jonathan Chan, Jessica Shi, and Stephanie Weirich. 2024b. Internalizing Indistinguishability with Dependent Types. *Proc. ACM Program. Lang.* 8, POPL, Article 44 (Jan. 2024), 28 pages. doi:10.1145/3632886
- Yiyun Liu, Jonathan Chan, and Stephanie Weirich. 2024a. *Artifact associated with Consistency of a Dependent Calculus of Indistinguishability*. University of Pennsylvania. doi:10.5281/zenodo.13930551
- Yiyun Liu, Jonathan Chan, and Stephanie Weirich. 2025. Consistency of a Dependent Calculus of Indistinguishability. *Proc. ACM Program. Lang.* 9, POPL (Jan. 2025), 27 pages. doi:10.1145
- Yiyun Liu and Stephanie Weirich. 2023. Dependently-Typed Programming with Logical Equality Reflection. *Proc. ACM Program. Lang.* 7, ICFP, Article 210 (Aug. 2023), 37 pages. doi:10.1145/3607852
- Per Martin-Löf. 1972. An intuitionistic theory of types.
- The mathlib Community. 2020. The lean mathematical library. In *Proceedings of the 9th ACM SIGPLAN International Conference on Certified Programs and Proofs* (New Orleans, LA, USA) (*CPP 2020*). Association for Computing Machinery, New York, NY, USA, 367–381. doi:10.1145/3372885.3373824
- Ulf Norell. 2007. *Towards a practical programming language based on dependent type theory*. Ph.D. Dissertation. Chalmers University of Technology and Göteborg University, Göteborg, Sweden. <https://research.chalmers.se/en/publication/46311>
- Frank Pfenning and Christine Paulin-Mohring. 1990. Inductively defined types in the Calculus of Constructions. In *Mathematical Foundations of Programming Semantics*, M. Main, A. Melton, M. Mislove, and D. Schmidt (Eds.). Vol. 442. Springer-Verlag, Berlin/Heidelberg, Germany, 209–228. doi:10.1007/BFb0040259
- Loïc Pujet and Nicolas Tabareau. 2022. Observational equality: now for good. *Proc. ACM Program. Lang.* 6, POPL, Article 32 (Jan. 2022), 27 pages. doi:10.1145/3498693
- Loïc Pujet and Nicolas Tabareau. 2023. Impredicative Observational Equality. *Proc. ACM Program. Lang.* 7, POPL, Article 74 (Jan. 2023), 26 pages. doi:10.1145/3571739
- Loïc Pujet and Nicolas Tabareau. 2024. Observational Equality Meets CIC. In *Programming Languages and Systems*, Stephanie Weirich (Ed.), Vol. 14576. Springer Nature Switzerland, Cham, 275–301. doi:10.1007/978-3-031-57262-3_12
- Aseem Rastogi, Guido Martínez, Aymeric Fromherz, Tahina Ramananandro, and Nikhil Swamy. 2021. Programming and Proving with Indexed Effects. <https://fstar-lang.org/papers/indexedeffects/indexedeffects.pdf>

- Vilhelm Sjöberg. 2015. *A Dependently Typed Language with Nontermination*. Ph.D. Dissertation. University of Pennsylvania. <https://repository.upenn.edu/handle/20.500.14332/27918>
- Thomas Streicher. 1993. *Investigations into intensional type theory*. Ph.D. Dissertation. Ludwig Maximilian Universität, Munich, Germany. <https://www2.mathematik.tu-darmstadt.de/~streicher/HabilStreicher.pdf>
- Nikhil Swamy, Cătălin Hrițcu, Chantal Keller, Aseem Rastogi, Antoine Delignat-Lavaud, Simon Forest, Karthikeyan Bhargavan, Cédric Fournet, Pierre-Yves Strub, Markulf Kohlweiss, Jean-Karim Zinzindohoue, and Santiago Zanella-Béguelin. 2016. Dependent types and multi-monadic effects in F*. *ACM SIGPLAN Notices* 51, 1 (Jan. 2016), 256–270. doi:10.1145/2914770.2837655
- William Walker Tait. 1967. Intensional interpretations of functionals of finite type I. *Journal of Symbolic Logic* 32, 2 (1967), 198–212. doi:10.2307/2271658
- The Univalent Foundations Program. 2013. *Homotopy Type Theory: Univalent Foundations of Mathematics*. <https://homotopytypetheory.org/book>, Institute for Advanced Study.
- Andrea Vezzosi, Anders Mörtberg, and Andreas Abel. 2019. Cubical Agda: a dependently typed programming language with univalence and higher inductive types. *Proc. ACM Program. Lang.* 3, ICFP, Article 87 (July 2019), 29 pages. doi:10.1145/3341691
- Stephanie Weirich. 2022. Implementing Dependent Types in pi-forall. doi:10.48550/arxiv.2207.02129 Lecture notes for the Oregon Programming Languages Summer School.
- Stephanie Weirich, Antoine Voizard, Pedro Henrique Azevedo de Amorim, and Richard A. Eisenberg. 2017. A specification for dependent types in Haskell. *Proc. ACM Program. Lang.* 1, ICFP, Article 31 (Aug. 2017), 29 pages. doi:10.1145/3110275
- Théo Winterhalter. 2024. Dependent Ghosts Have a Reflection for Free. *Proc. ACM Program. Lang.* 8, ICFP, Article 258 (Aug. 2024), 29 pages. doi:10.1145/3674647