

Coding Sequences as Executable Functions: A Type-Theoretic Framework for the Central Dogma

Matthew Long

The YonedaAI Collaboration

Chicago, IL

March 31, 2026

bio.PL

Abstract

We present a comprehensive type-theoretic framework that formalizes the central dogma of molecular biology—the flow of genetic information from DNA through mRNA to protein—as a typed compilation pipeline. Under this formalization, a protein-coding DNA sequence constitutes a *source program* written in a four-letter alphabet $\{A, C, G, T\}$; transcription to mRNA yields an *intermediate representation* over $\{A, C, G, U\}$; and ribosomal translation to a polypeptide chain constitutes *compilation to executable binary* over the twenty-letter amino acid alphabet. We introduce the parametric type $\text{ProteinCode}\langle T \rangle$, which encapsulates the entire compilation pipeline from gene to gene product, and we establish a formal codon type system in which the standard genetic code’s $64 \rightarrow 20 + \text{Stop}$ surjection provides error-correcting redundancy analogous to channel coding. Point mutations are classified as identity refactorings (silent), type-compatible substitutions (missense), type errors (nonsense), and parse errors (frameshift). Alternative splicing is formalized via dependent types indexed by cellular context, and post-translational modifications are modeled as endofunctors on the category of protein structures. We prove three principal theorems: a *Compilation Correspondence Theorem* establishing the functoriality of the central dogma, a *Codon Redundancy Theorem* quantifying error-correction capacity, and a *Splicing Sheaf Theorem* characterizing the presheaf structure of alternative splicing. A reference implementation in Haskell demonstrates the framework’s computational content.

Keywords: type theory, central dogma, genetic code, compilation pipeline, codon type system, alternative splicing, dependent types, category theory, DNA-Lang

Contents

1	Introduction	4
2	The Central Dogma as Compilation	5
2.1	Alphabets and Sequence Spaces	5
2.2	Transcription: Source to IR	6
2.3	Translation: IR to Binary	6
2.4	The Ribosome as Virtual Machine	7
2.5	Pre-mRNA Processing as Optimization	8
3	The Codon Type System	8
3.1	The Codon Space	9
3.2	Degeneracy Structure	9
3.3	Wobble Position and Third-Position Tolerance	10
3.4	Reading Frames	10
3.5	The Codon Table as a Graph Coloring	10
3.6	Information-Theoretic Analysis	11
4	ProteinCode$\langle T \rangle$: A Parametric Type for Gene Expression	11
4.1	The Type Definition	11
4.2	Type Rules for Transcription and Translation	12
4.3	The Expression Judgment	12
4.4	ProteinCode as a Functor	13
4.5	The Compilation Commutative Diagram	14
5	Mutations as Type Errors	14
5.1	Classification of Point Mutations	14
5.2	The Software Engineering Correspondence	14
5.3	Formal Mutation Algebra	15
5.4	Sickle Cell Disease as a Missense Example	15
5.5	Nonsense Mutations and Premature Termination	16
5.6	Frameshifts as Parsing Failures	16
5.7	Mutation Severity Ordering	17
6	Alternative Splicing as Dependent Types	17
6.1	The Splicing Function	18
6.2	Dependent Typing of Gene Products	18
6.3	The Splicing Presheaf	19
6.4	The Dependent Product and Sum Types	19
6.5	Combinatorial Complexity	20

7	Post-Translational Modification as Runtime Decoration	20
7.1	PTM as Endofunctors	20
7.2	Composition of PTMs	21
7.3	Natural Transformations Between PTM States	21
7.4	The PTM Enrichment Diagram	22
7.5	The Histone Code as a PTM Algebra	22
8	Formal Theorems	23
8.1	The Compilation Correspondence Theorem	23
8.2	The Codon Redundancy Error Correction Theorem	24
8.3	The Splicing Sheaf Theorem	25
9	The Haskell Implementation	26
9.1	Module Architecture	27
9.2	Core Types: ProteinCode.hs	27
9.3	The Codon Table: CodonTable.hs	29
9.4	Transcription: Transcription.hs	30
9.5	Translation: Translation.hs	31
9.6	Demonstration: Main.hs	32
10	Discussion and Conclusion	33
10.1	Summary of Contributions	33
10.2	Relationship to Prior Work	34
10.3	The Compiler Architecture of the Cell	34
10.4	Implications for Synthetic Biology	35
10.5	Limitations and Future Work	35
10.6	Conclusion	36
A	Complete Codon Table	39
B	Categorical Glossary	39
C	Notation Index	41

1 Introduction

The central dogma of molecular biology, first articulated by Crick [1, 2], describes the unidirectional flow of sequence information from nucleic acid to protein. In its modern form, the dogma asserts that DNA is transcribed into messenger RNA (mRNA), which is in turn translated by the ribosome into a polypeptide chain that folds into a functional protein. This information flow—DNA $\rightarrow$ mRNA $\rightarrow$ Protein—is the most fundamental compilation pipeline in biology.

The analogy between genetic information processing and computation has been explored since the earliest days of molecular biology. Schrödinger’s prescient notion of a “hereditary code-script” [4] anticipated both the discovery of the double helix and the development of information theory. Subsequent work by Gamow, Crick, and Nirenberg on the genetic code [3] explicitly framed codon assignment as a coding-theoretic problem. More recently, researchers have drawn connections between molecular biology and formal language theory [5], noting that the structure of gene expression bears striking parallels to the architecture of compilers and interpreters.

In this paper, we develop these parallels into a rigorous mathematical framework. Our central insight is that the compilation pipeline metaphor is not merely a pedagogical analogy but admits precise formalization in the language of type theory and category theory. Specifically:

- (i) **DNA as source code.** A protein-coding gene is a source program written in the four-letter alphabet $\Sigma_{\text{DNA}} = \{A, C, G, T\}$. The gene contains both coding regions (exons) and non-coding regions (introns), analogous to a source file containing both executable statements and comments or preprocessor directives.
- (ii) **mRNA as intermediate representation.** Transcription produces an intermediate representation in the alphabet $\Sigma_{\text{RNA}} = \{A, C, G, U\}$, where the substitution $T \mapsto U$ constitutes a well-typed alphabet transformation. Splicing of the pre-mRNA to remove introns corresponds to an optimization pass that eliminates dead code and produces a streamlined IR.
- (iii) **Protein as executable binary.** Translation by the ribosome compiles the mRNA into a polypeptide—an executable program over the twenty-letter amino acid alphabet Σ_{AA} . The ribosome functions as a virtual machine that reads the IR in triplet codons and emits amino acid instructions.
- (iv) **The codon table as a type system.** The standard genetic code, which maps 64 codons to 20 amino acids plus a stop signal, constitutes a surjective

type assignment. The redundancy inherent in this mapping provides error-correcting properties: many single-nucleotide substitutions are absorbed by the type system as silent mutations.

The remainder of this paper is organized as follows. Section 2 formalizes the central dogma as a two-pass compilation pipeline. Section 3 develops the codon type system and analyzes its error-correcting properties. Section 4 introduces the parametric type $\text{ProteinCode}\langle T \rangle$ that unifies the pipeline. Section 5 classifies mutations in type-theoretic terms. Section 6 formalizes alternative splicing via dependent types. Section 7 models post-translational modifications as endofunctors. Section 8 states and proves the principal theorems. Section 9 describes the Haskell reference implementation. Section 10 concludes with discussion and future directions.

2 The Central Dogma as Compilation

We begin by establishing the categorical framework for the central dogma viewed as a multi-pass compiler.

2.1 Alphabets and Sequence Spaces

Definition 2.1 (Biological Alphabets). We define three alphabets corresponding to the three stages of the central dogma:

$$\Sigma_{\text{DNA}} = \{A, C, G, T\}, \quad (1)$$

$$\Sigma_{\text{RNA}} = \{A, C, G, U\}, \quad (2)$$

$$\begin{aligned} \Sigma_{\text{AA}} = \{ & \text{Ala, Arg, Asn, Asp, Cys, Gln, Glu, Gly, His, Ile,} \\ & \text{Leu, Lys, Met, Phe, Pro, Ser, Thr, Trp, Tyr, Val, Stop} \}. \end{aligned} \quad (3)$$

The free monoids Σ_{DNA}^* , Σ_{RNA}^* , and Σ_{AA}^* are the corresponding sequence spaces.

Definition 2.2 (Compilation Pipeline). The central dogma is a composition of two functors in the category **Seq** of biological sequences:

$$\mathcal{T}: \Sigma_{\text{DNA}}^* \rightarrow \Sigma_{\text{RNA}}^*, \quad \mathcal{R}: \Sigma_{\text{RNA}}^* \rightarrow \Sigma_{\text{AA}}^*,$$

where $\mathcal{T}$ is the *transcription functor* and $\mathcal{R}$ is the *ribosomal translation functor*. The composite

$$\mathcal{E} = \mathcal{R} \circ \mathcal{T}: \Sigma_{\text{DNA}}^* \rightarrow \Sigma_{\text{AA}}^*$$

is the *expression functor* mapping genes to their protein products.

2.2 Transcription: Source to IR

The transcription functor $\mathcal{T}$ operates as a character-level map on the sense strand:

Definition 2.3 (Transcription Map). The transcription map $\tau: \Sigma_{\text{DNA}} \rightarrow \Sigma_{\text{RNA}}$ is defined by:

$$\tau(A) = A, \quad \tau(C) = C, \quad \tau(G) = G, \quad \tau(T) = U.$$

This is an isomorphism of alphabets (bijection on the four-element sets) and extends to a monoid homomorphism $\mathcal{T} = \tau^*: \Sigma_{\text{DNA}}^* \rightarrow \Sigma_{\text{RNA}}^*$ by the universal property of the free monoid.

Proposition 2.4 (Transcription is Invertible). *The transcription map τ is a bijection with inverse τ^{-1} given by $U \mapsto T$ and identity on $\{A, C, G\}$. Consequently, $\mathcal{T}$ is an isomorphism of free monoids.*

Proof. By direct verification: $\tau^{-1} \circ \tau = \text{id}_{\Sigma_{\text{DNA}}}$ and $\tau \circ \tau^{-1} = \text{id}_{\Sigma_{\text{RNA}}}$. The extension to free monoids preserves this property since free monoid homomorphisms are determined by their action on generators. $\square$

Remark 2.5. The invertibility of transcription corresponds to the compiler-theoretic principle that the source-to-IR transformation should be *information-preserving*: no semantic content is lost in the translation from DNA to mRNA. The biological reality of reverse transcription (by retroviruses) confirms this invertibility empirically.

2.3 Translation: IR to Binary

Translation, unlike transcription, is a *lossy* compilation pass. The ribosome reads the mRNA in non-overlapping triplets (codons) and maps each codon to an amino acid via the genetic code.

Definition 2.6 (Codon Extraction). The *codon extraction* function in reading frame $f \in \{0, 1, 2\}$ is:

$$\text{codons}_f: \Sigma_{\text{RNA}}^* \rightarrow (\Sigma_{\text{RNA}}^3)^*, \quad (r_1 r_2 \cdots r_n) \mapsto ((r_{f+1}, r_{f+2}, r_{f+3}), (r_{f+4}, r_{f+5}, r_{f+6}), \dots).$$

The standard reading frame is $f = 0$ when initiated from a start codon (AUG).

Definition 2.7 (The Genetic Code). The *standard genetic code* is the surjection:

$$\gamma: \Sigma_{\text{RNA}}^3 \twoheadrightarrow \Sigma_{\text{AA}} \cup \{\text{Stop}\},$$

mapping each of the $|\Sigma_{\text{RNA}}^3| = 4^3 = 64$ codons to one of 20 amino acids or a stop signal.

Definition 2.8 (Translation Functor). The *translation functor* $\mathcal{R}$ acts on an mRNA sequence $m \in \Sigma_{\text{RNA}}^*$ by:

- (a) locating the first AUG codon (start signal),
- (b) extracting codons in reading frame $f = 0$ from that position,
- (c) applying γ to each codon,
- (d) terminating at the first Stop signal.

Formally, if $m = m_1 \cdots m_k \cdot \text{AUG} \cdot c_1 c_2 c_3 \cdots$ where $c_i \in \Sigma_{\text{RNA}}$, then

$$\mathcal{R}(m) = \gamma(c_1 c_2 c_3) \cdot \gamma(c_4 c_5 c_6) \cdots \gamma(c_{3j-2} c_{3j-1} c_{3j})$$

where j is the smallest index such that $\gamma(c_{3(j+1)-2} c_{3(j+1)-1} c_{3(j+1)}) = \text{Stop}$.

2.4 The Ribosome as Virtual Machine

The ribosome can be modeled as a state machine—specifically, a deterministic pushdown automaton—that processes the mRNA tape:

Definition 2.9 (Ribosomal VM). The ribosomal virtual machine is a tuple $\mathcal{V} = (Q, \Sigma_{\text{RNA}}, \Sigma_{\text{AA}}, \delta, q_0, F)$ where:

- $Q = \{q_{\text{scan}}, q_{\text{init}}, q_{\text{elong}}, q_{\text{term}}\}$ is the state set,
- δ is the transition function consuming three input symbols per step,
- $q_0 = q_{\text{scan}}$ is the initial state (scanning for AUG),
- $F = \{q_{\text{term}}\}$ is the accepting state (stop codon reached).

The execution phases correspond to:

1. **Initiation** ($q_{\text{scan}} \rightarrow q_{\text{init}}$): Scan for the AUG start codon. Analogous to the bootloader locating the entry point.
2. **Elongation** ($q_{\text{init}} \rightarrow q_{\text{elong}}$): Iteratively read codons and emit amino acids. Analogous to the fetch-decode-execute cycle.
3. **Termination** ($q_{\text{elong}} \rightarrow q_{\text{term}}$): Encounter a stop codon and release the completed polypeptide. Analogous to program halt.

The following commutative diagram summarizes the compilation pipeline:

$$\begin{array}{ccc}
 \Sigma_{\text{DNA}}^* & \xrightarrow[\text{Transcription}]{\mathcal{T}} & \Sigma_{\text{RNA}}^* \\
 & \searrow \mathcal{E} = \mathcal{R} \circ \mathcal{T} & \downarrow \mathcal{R} \\
 & & \Sigma_{\text{AA}}^*
 \end{array} \tag{4}$$

2.5 Pre-mRNA Processing as Optimization

Between transcription and translation, the pre-mRNA undergoes several processing steps that have direct analogues in compiler optimization:

Definition 2.10 (Pre-mRNA Processing Pipeline). The pre-mRNA processing pipeline consists of three operations:

- (a) **5' capping**: Addition of a 7-methylguanosine cap to the 5' end. Analogous to adding a file header or magic number that identifies the format.
- (b) **3' polyadenylation**: Addition of a poly(A) tail to the 3' end. Analogous to adding a checksum or sentinel value that marks the end of valid data.
- (c) **Splicing**: Removal of introns and joining of exons. Analogous to dead code elimination and function inlining—the most significant optimization pass.

The splicing operation is formalized as a morphism in the category of sequences:

Definition 2.11 (Splicing Morphism). Given a gene g with exon set $\{E_1, \dots, E_n\}$ and intron set $\{I_1, \dots, I_{n-1}\}$ (where the pre-mRNA has the form $E_1 I_1 E_2 I_2 \dots E_n$), the splicing morphism is:

$$\text{Splice}: \Sigma_{\text{RNA}}^* \rightarrow \Sigma_{\text{RNA}}^*, \quad E_1 I_1 E_2 I_2 \dots E_n \mapsto E_1 E_2 \dots E_n.$$

This is a retraction (left inverse) of the inclusion $E_1 E_2 \dots E_n \hookrightarrow E_1 I_1 E_2 I_2 \dots E_n$.

3 The Codon Type System

The genetic code constitutes a type system in which 64 syntactic forms (codons) are assigned to 21 semantic types (20 amino acids plus Stop). We now formalize this type assignment and analyze its error-correcting properties.

3.1 The Codon Space

Definition 3.1 (Codon Space). The *codon space* is the set $\mathcal{C} = \Sigma_{\text{RNA}}^3 = \{A, C, G, U\}^3$ of all 64 ordered triples of RNA nucleotides. We equip $\mathcal{C}$ with the Hamming metric:

$$d_H(c, c') = |\{i \in \{1, 2, 3\} : c_i \neq c'_i\}|.$$

The Hamming ball of radius r around a codon c is $B_r(c) = \{c' \in \mathcal{C} : d_H(c, c') \leq r\}$.

Proposition 3.2 (Hamming Ball Sizes). *In the codon space $\mathcal{C}$:*

$$|B_0(c)| = 1, \tag{5}$$

$$|B_1(c)| = 1 + 3 \times 3 = 10, \tag{6}$$

$$|B_2(c)| = 1 + 9 + 3 \times 9 = 37, \tag{7}$$

$$|B_3(c)| = 64. \tag{8}$$

Proof. $B_0(c)$ contains only c . For $B_1(c)$, each of 3 positions admits 3 alternative nucleotides, giving 9 neighbors plus c itself. For $B_2(c)$, we add $\binom{3}{2} \times 3^2 = 27$ codons differing in exactly 2 positions. For $B_3(c)$, $\binom{3}{3} \times 3^3 = 27$ codons differ in all 3 positions, and $1 + 9 + 27 + 27 = 64$. $\square$

Definition 3.3 (Type Assignment). The *codon type assignment* is the function $\gamma: \mathcal{C} \rightarrow \mathcal{A}$ where $\mathcal{A} = \Sigma_{\text{AA}} \cup \{\text{Stop}\}$, as defined in Theorem 2.7. For each amino acid $a \in \mathcal{A}$, the *codon class* or *synonymous codon set* is:

$$[a] = \gamma^{-1}(a) = \{c \in \mathcal{C} : \gamma(c) = a\}.$$

3.2 Degeneracy Structure

The codon type system exhibits a characteristic pattern of degeneracy:

Proposition 3.4 (Degeneracy Distribution). *The sizes of the codon classes under the standard genetic code are:*

<i>Degeneracy</i>	<i>Count</i>	<i>Amino Acids / Signals</i>
6-fold	3	Leu, Ser, Arg
4-fold	5	Val, Pro, Thr, Ala, Gly
3-fold	1	Ile
2-fold	9	Phe, Tyr, His, Gln, Asn, Lys, Asp, Glu, Cys
1-fold	2	Met (AUG), Trp (UGG)
3-fold (stop)	1	Stop (UAA, UAG, UGA)

Verification: $3 \times 6 + 5 \times 4 + 1 \times 3 + 9 \times 2 + 2 \times 1 + 1 \times 3 = 18 + 20 + 3 + 18 + 2 + 3 = 64$.

3.3 Wobble Position and Third-Position Tolerance

Definition 3.5 (Wobble Equivalence). Define the *wobble equivalence relation* $\sim_w$ on $\mathcal{C}$ by

$$(c_1, c_2, c_3) \sim_w (c'_1, c'_2, c'_3) \iff c_1 = c'_1 \text{ and } c_2 = c'_2.$$

This partitions $\mathcal{C}$ into 16 wobble classes, each of size 4.

Proposition 3.6 (Third-Position Degeneracy). *For 8 of the 16 wobble classes, all four members encode the same amino acid (the “four-fold degenerate” sites). For the remaining 8 classes, the four codons split into two pairs encoding distinct amino acids (or amino acid and stop). Consequently, a point mutation at the third codon position has at most a 50% chance of altering the amino acid, and for four-fold degenerate sites, the probability is 0%.*

Proof. By exhaustive enumeration of the codon table. The four-fold degenerate wobble classes are: UCx (Ser), CCx (Pro), ACx (Thr), GCx (Ala), CUx (Leu), GUx (Val), CGx (Arg), and GGx (Gly). For each, $\gamma(c_1, c_2, x)$ is constant over $x \in \{A, C, G, U\}$. $\square$

3.4 Reading Frames

Definition 3.7 (Reading Frame). For an mRNA sequence $m = m_1 m_2 \cdots m_n$, the reading frame $f \in \{0, 1, 2\}$ determines the phase of codon extraction. The three forward reading frames yield potentially distinct protein products:

$$\pi_f(m) = \mathcal{R}_f(m) = \gamma(m_{f+1}, m_{f+2}, m_{f+3}) \cdot \gamma(m_{f+4}, m_{f+5}, m_{f+6}) \cdots$$

In the biological context, the reading frame is established by the position of the AUG initiation codon.

Remark 3.8. The existence of three reading frames per strand (and three on the complementary strand, for a total of six) means that any given stretch of nucleotide sequence potentially encodes six overlapping protein products. This is analogous to polyglot programs that are simultaneously valid in multiple programming languages, depending on which parser (reading frame) is applied.

3.5 The Codon Table as a Graph Coloring

We can view the codon type assignment as a graph coloring problem.

Definition 3.9 (Hamming Graph of Codons). Let $G_{\mathcal{C}} = (\mathcal{C}, E)$ be the graph where codons c, c' are adjacent if $d_H(c, c') = 1$ (i.e., they differ in exactly one position).

Each codon has exactly $3 \times 3 = 9$ neighbors (three positions, each with three alternative nucleotides), so G_C is a 9-regular graph on 64 vertices.

Proposition 3.10 (Silent Mutation Clustering). *The codon type assignment γ exhibits clustering: codons encoding the same amino acid tend to form connected subgraphs in G_C , particularly through third-position variation. The average number of synonymous neighbors per codon is approximately 4.2, compared to the expected value of $9 \times (|\gamma^{-1}(\gamma(c))| - 1)/63$ under a random assignment.*

This clustering is the mathematical signature of error tolerance: the genetic code is optimized to minimize the phenotypic impact of single-nucleotide substitutions.

3.6 Information-Theoretic Analysis

Proposition 3.11 (Information Content of the Genetic Code). *The information content of the codon-to-amino-acid mapping can be quantified as follows:*

- (a) *The raw information capacity of a codon is $\log_2 64 = 6$ bits.*
- (b) *The information needed to specify one of 20 amino acids (ignoring Stop) is $\log_2 20 \approx 4.32$ bits.*
- (c) *The redundancy is $6 - 4.32 = 1.68$ bits per codon, representing the error-correcting overhead.*
- (d) *The coding efficiency is $4.32/6 \approx 72\%$, with 28% devoted to error tolerance.*

4 ProteinCode $\langle T \rangle$: A Parametric Type for Gene Expression

We now introduce the central type-theoretic construction of this paper: the parametric type $\text{ProteinCode}\langle T \rangle$, which encapsulates the complete compilation pipeline from a gene to its product.

4.1 The Type Definition

Definition 4.1 ($\text{ProteinCode}\langle T \rangle$ Type). Let T be a type representing a gene product (typically a protein, but potentially an RNA molecule for non-coding genes).

The type $\text{ProteinCode}\langle T \rangle$ is a record type:

$$\text{ProteinCode}\langle T \rangle = \left\{ \begin{array}{ll} \text{gene} & : \Sigma_{\text{DNA}}^* \\ \text{product} & : T \\ \text{exons} & : \text{List}(\text{Exon}) \\ \text{name} & : \text{String} \end{array} \right\}$$

where **Exon** records an exon identifier, genomic coordinates, and its nucleotide sequence.

Definition 4.2 (Exon and Gene Structure). An **Exon** is a tuple (i, s, e, σ) where $i \in \mathbb{N}$ is the exon index, $s, e \in \mathbb{N}$ are the start and end coordinates, and $\sigma \in \Sigma_{\text{DNA}}^*$ is the exon sequence. A **GeneStructure** consists of:

$$\text{GeneStructure} = \left\{ \begin{array}{ll} \text{name} & : \text{String} \\ \text{exons} & : \text{List}(\text{Exon}) \\ \text{introns} & : \text{List}(\text{Intron}) \end{array} \right\}$$

where introns are the non-coding intervening sequences between exons.

4.2 Type Rules for Transcription and Translation

The transcription pass can be formalized as a type judgment:

Definition 4.3 (Transcription Judgment). The transcription typing rule is:

$$\frac{\Gamma \vdash g : \Sigma_{\text{DNA}}^* \quad \tau : \Sigma_{\text{DNA}} \xrightarrow{\sim} \Sigma_{\text{RNA}}}{\Gamma \vdash \mathcal{T}(g) : \Sigma_{\text{RNA}}^*} \quad (\text{T-TRANSCRIBE})$$

This rule states: given a well-typed DNA sequence g and the alphabet isomorphism τ , the transcribed sequence $\mathcal{T}(g)$ is a well-typed RNA sequence.

Definition 4.4 (Translation Judgment). The translation typing rule is:

$$\frac{\Gamma \vdash m : \Sigma_{\text{RNA}}^* \quad \text{hasStart}(m) \quad \gamma : \mathcal{C} \twoheadrightarrow \mathcal{A}}{\Gamma \vdash \mathcal{R}(m) : \Sigma_{\text{AA}}^*} \quad (\text{T-TRANSLATE})$$

The precondition $\text{hasStart}(m)$ asserts that the mRNA contains at least one AUG codon, establishing the reading frame.

4.3 The Expression Judgment

Composing the two judgments, we obtain the complete expression judgment:

Definition 4.5 (Expression Judgment). The gene expression typing rule is:

$$\frac{\Gamma \vdash g : \Sigma_{\text{DNA}}^* \quad \Gamma \vdash \mathcal{T}(g) : \Sigma_{\text{RNA}}^* \quad \text{hasStart}(\mathcal{T}(g))}{\Gamma \vdash \mathcal{E}(g) : \Sigma_{\text{AA}}^*} \quad (\text{T-EXPRESS})$$

Proposition 4.6 (Type Safety of Expression). *If a gene g is well-formed (i.e., $g \in \Sigma_{\text{DNA}}^*$, contains at least one ATG codon, and has an in-frame stop codon downstream), then the expression judgment $\Gamma \vdash \mathcal{E}(g) : \Sigma_{\text{AA}}^*$ is derivable. That is: well-formed genes always produce well-typed proteins.*

Proof. Well-formedness of g ensures $g \in \Sigma_{\text{DNA}}^*$, so $\mathcal{T}(g) \in \Sigma_{\text{RNA}}^*$ by Theorem 4.3. The presence of ATG in g implies AUG in $\mathcal{T}(g)$ (since $\tau(T) = U$ and τ fixes A, G), so $\text{hasStart}(\mathcal{T}(g))$ holds. An in-frame stop codon ensures termination, producing a finite element of Σ_{AA}^* . $\square$

4.4 ProteinCode as a Functor

Proposition 4.7 ($\text{ProteinCode}\langle T \rangle$ is Functorial). *The type constructor $\text{ProteinCode}\langle - \rangle$ defines a functor from the category of gene products to the category of annotated gene records. Given a morphism $f : T \rightarrow T'$ of gene products (e.g., a protein-protein interaction or a post-translational modification), we obtain:*

$$\text{ProteinCode}\langle f \rangle : \text{ProteinCode}\langle T \rangle \rightarrow \text{ProteinCode}\langle T' \rangle$$

*defined by applying f to the **product** field while preserving the gene, exon, and name fields.*

Proof. We verify the functor laws. Identity: $\text{ProteinCode}\langle \text{id}_T \rangle$ preserves all fields, acting as $\text{id}_{\text{ProteinCode}\langle T \rangle}$. Composition: $\text{ProteinCode}\langle g \circ f \rangle = \text{ProteinCode}\langle g \rangle \circ \text{ProteinCode}\langle f \rangle$ follows from the fact that only the product field is transformed, and function composition is associative. $\square$

Corollary 4.8. *The assignment $T \mapsto \text{ProteinCode}\langle T \rangle$ is an endofunctor on **Type** (the category of types and functions). This endofunctor admits a comonad structure $(\text{ProteinCode}, \varepsilon, \delta)$ where:*

- $\varepsilon : \text{ProteinCode}\langle T \rangle \rightarrow T$ extracts the product (the **pcProduct** field),
- $\delta : \text{ProteinCode}\langle T \rangle \rightarrow \text{ProteinCode}\langle \text{ProteinCode}\langle T \rangle \rangle$ duplicates the annotation by wrapping the entire record as the new product.

The comonad laws $\varepsilon \circ \delta = \text{id}$ and $\text{ProteinCode}\langle \varepsilon \rangle \circ \delta = \text{id}$ follow from the definitions.

4.5 The Compilation Commutative Diagram

The full type-theoretic structure of the central dogma is captured by the following commutative diagram in **Type**:

5 Mutations as Type Errors

Point mutations—substitutions of a single nucleotide—can be classified by their effect on the type system. This classification directly parallels the taxonomy of code changes in software engineering.

5.1 Classification of Point Mutations

Definition 5.1 (Mutation Type). Let $c = (c_1, c_2, c_3) \in \mathcal{C}$ be a codon and $c' = (c'_1, c'_2, c'_3) \in \mathcal{C}$ a codon obtained by a single-nucleotide substitution ($d_H(c, c') = 1$). The mutation is classified as:

- (i) **Silent** if $\gamma(c) = \gamma(c')$ and $\gamma(c) \neq \text{Stop}$. The amino acid is unchanged. This is *identity refactoring*: the source code has changed, but the compiled binary is identical.
- (ii) **Missense** if $\gamma(c) \neq \gamma(c')$ and both $\gamma(c), \gamma(c') \in \Sigma_{AA}$ (neither is Stop). A different amino acid is substituted. This is a *type-compatible substitution*: the program still type-checks, but its behavior may differ.
- (iii) **Nonsense** if $\gamma(c) \in \Sigma_{AA}$ and $\gamma(c') = \text{Stop}$. A premature termination signal is introduced. This is a *type error*: the compilation aborts before producing a complete output.
- (iv) **Frameshift** (for insertions/deletions) if the mutation shifts the reading frame. This is a *parse error*: the lexer's tokenization boundary is corrupted, and all downstream codons are misread.

5.2 The Software Engineering Correspondence

The following table makes the correspondence precise:

Table 1: Mutation types and their software engineering analogues.

Biological Term	Type-Theoretic Term	SE Analogue	Severity
Silent mutation	Identity refactoring	Rename local variable	None
Missense mutation	Compatible substitution	Change function body	Variable
Nonsense mutation	Type error (abort)	Missing return statement	Critical
Frameshift	Parse error	Missing delimiter	Catastrophic

5.3 Formal Mutation Algebra

Definition 5.2 (Mutation Operator). A *point mutation operator* at position $i \in \{1, 2, 3\}$ and target nucleotide $n \in \Sigma_{\text{RNA}}$ is:

$$\mu_{i,n}: \mathcal{C} \rightarrow \mathcal{C}, \quad \mu_{i,n}(c_1, c_2, c_3) = (c'_1, c'_2, c'_3) \text{ where } c'_j = \begin{cases} n & \text{if } j = i, \\ c_j & \text{otherwise.} \end{cases}$$

Proposition 5.3 (Mutation Operators are Idempotent). *Each mutation operator $\mu_{i,n}$ satisfies $\mu_{i,n} \circ \mu_{i,n} = \mu_{i,n}$ (idempotency). The set of all mutation operators at a fixed position i generates a monoid isomorphic to the transformation monoid on $\{A, C, G, U\}$ acting on the i -th coordinate.*

Proof. Idempotency follows immediately: applying $\mu_{i,n}$ twice sets coordinate i to n both times. The monoid structure follows from the fact that composing two mutations at the same position yields another mutation at that position: $\mu_{i,m} \circ \mu_{i,n} = \mu_{i,m}$ (the second application overwrites the first). $\square$

Definition 5.4 (Mutation Classification Functor). Define the functor $\mathcal{M}: \mathbf{Codon} \rightarrow \mathbf{MutType}$ from the category of codons (with mutation operators as morphisms) to the discrete category of mutation types $\{\text{Silent}, \text{Missense}, \text{Nonsense}, \text{Frameshift}\}$ by:

$$\mathcal{M}(\mu_{i,n})(c) = \begin{cases} \text{Silent} & \text{if } \gamma(\mu_{i,n}(c)) = \gamma(c), \\ \text{Nonsense} & \text{if } \gamma(\mu_{i,n}(c)) = \text{Stop} \text{ and } \gamma(c) \neq \text{Stop}, \\ \text{Missense} & \text{otherwise.} \end{cases}$$

5.4 Sickle Cell Disease as a Missense Example

Example 5.5 (Sickle Cell Mutation). The sickle cell mutation replaces glutamic acid (Glu) with valine (Val) at position 6 of the β -globin chain. At the codon level:

$$\text{GAG} \xrightarrow{\mu_{2,U}} \text{GUG}, \quad \gamma(\text{GAG}) = \text{Glu}, \quad \gamma(\text{GUG}) = \text{Val}.$$

This is a missense mutation: the type changes from Glu to Val. The substitution is type-compatible (both are amino acids), but the behavioral change is dramatic—the valine residue causes hemoglobin polymerization under low-oxygen conditions, producing the characteristic sickle shape.

In software terms, this is analogous to changing a single function parameter type from a hydrophilic interface to a hydrophobic one: the program still compiles, but the runtime behavior is pathological.

5.5 Nonsense Mutations and Premature Termination

Proposition 5.6 (Nonsense Density). *Of the $64 \times 9 = 576$ possible single-nucleotide substitutions across all codons, exactly those that produce a stop codon from a sense codon are nonsense mutations. Excluding stop-to-stop transitions, there are 23 distinct sense-codon-to-stop-codon transitions, yielding a nonsense mutation rate of approximately $23/549 \approx 4.2\%$ of all non-trivial substitutions from sense codons.*

Example 5.7 (Nonsense Mutation: Glutamine to Stop). The codon CAG (Gln) can mutate to UAG (Amber stop) via $\mu_{1,U}$:

$$\text{CAG} \xrightarrow{\mu_{1,U}} \text{UAG}, \quad \gamma(\text{CAG}) = \text{Gln}, \quad \gamma(\text{UAG}) = \text{Stop}.$$

This introduces a premature termination signal, truncating the protein. In type-theoretic terms, the function’s return type has changed from a value type to $\perp$ (the bottom type), aborting execution.

5.6 Frameshifts as Parsing Failures

Definition 5.8 (Frameshift Perturbation). An *insertion frameshift* at position k in an mRNA sequence $m = m_1 \cdots m_n$ produces:

$$m' = m_1 \cdots m_k \cdot n_{\text{ins}} \cdot m_{k+1} \cdots m_n$$

where $n_{\text{ins}} \in \Sigma_{\text{RNA}}$. For all positions $j > k$, the reading frame shifts by $+1 \pmod{3}$, causing every downstream codon to be misread.

Proposition 5.9 (Frameshift Catastrophe). *Let m be an mRNA with n codons downstream of a frameshift at position k . The expected number of codons until a stop codon is encountered in the shifted frame follows a geometric distribution with success probability $p = 3/64$. The expected run length before termination is:*

$$\mathbb{E}[\text{codons to stop}] = \frac{1-p}{p} = \frac{61}{3} \approx 20.3.$$

Therefore, frameshifts typically produce short, aberrant polypeptides—truncated programs that fail at runtime.

Proof. In the shifted frame, codons are essentially drawn from a uniform distribution over $\mathcal{C}$ (since the frame has no a priori relationship to the original coding sequence). The probability of any given codon being a stop codon is $3/64$. The number of non-stop codons before the first stop is geometrically distributed with parameter $p = 3/64$, giving the stated expectation. $\square$

5.7 Mutation Severity Ordering

Definition 5.10 (Severity Partial Order). We define a partial order on mutation types by their type-theoretic severity:

$$\text{Silent} \prec \text{Missense} \prec \text{Nonsense} \prec \text{Frameshift}.$$

This ordering reflects increasing disruption to the compilation pipeline: Silent preserves the output exactly, Missense changes a single instruction, Nonsense truncates the output, and Frameshift corrupts all downstream instructions.

Lemma 5.11 (Mutation Composition Monotonicity). *Let μ_1, μ_2 be two independent point mutations applied to a coding sequence. Then:*

$$\text{severity}(\mu_2 \circ \mu_1) \geq \max(\text{severity}(\mu_1), \text{severity}(\mu_2)).$$

That is, the combined effect of two mutations is at least as severe as either mutation alone.

Proof. If either mutation is a frameshift, the composition is a frameshift (or worse). If one is nonsense, the truncation is at least as early. If both are missense, the combined effect changes two amino acids, which is at least as severe as changing one. Silent mutations composed with anything yield the same or greater severity. $\square$

6 Alternative Splicing as Dependent Types

One of the most striking features of eukaryotic gene expression is *alternative splicing*: a single gene can produce multiple distinct proteins depending on which exons are included in the mature mRNA. This phenomenon maps naturally to the concept of *dependent types*—types that depend on a value (the cellular context).

6.1 The Splicing Function

Definition 6.1 (Splicing Context). A *splicing context* is a value $\sigma \in \mathcal{S}$ drawn from a context set:

$$\mathcal{S} = \{\text{Neural}, \text{Muscle}, \text{Epithelial}, \text{Hepatic}, \text{Default}, \dots\}$$

representing the tissue type, developmental stage, or physiological condition of the cell.

Definition 6.2 (Exon Selection Function). The *exon selection function* is:

$$\text{sel}: \text{List}(\text{Exon}) \times \mathcal{S} \rightarrow \text{List}(\text{Exon}),$$

which selects a subset of the gene's exons based on the splicing context. The mature mRNA is then:

$$\text{mRNA}(g, \sigma) = \bigoplus_{e \in \text{sel}(\text{exons}(g), \sigma)} \mathcal{T}(\text{seq}(e)),$$

where $\bigoplus$ denotes concatenation and $\text{seq}(e)$ is the nucleotide sequence of exon e .

6.2 Dependent Typing of Gene Products

Definition 6.3 (Context-Dependent Gene Product Type). The type of the protein product of a gene g in context σ is:

$$\text{Product}(g, \sigma) = \mathcal{R}(\text{mRNA}(g, \sigma)).$$

This is a *dependent type*: the type $\text{Product}(g, -)$ is a function $\mathcal{S} \rightarrow \mathbf{Type}$ that assigns a distinct protein type to each splicing context.

Example 6.4 (Tissue-Specific Isoforms). Consider a gene with four exons $\{E_1, E_2, E_3, E_4\}$. The splicing function might specify:

$$\begin{aligned} \text{sel}(_, \text{Default}) &= \{E_1, E_2, E_3, E_4\}, \\ \text{sel}(_, \text{Neural}) &= \{E_1, E_2, E_3, E_4\}, \\ \text{sel}(_, \text{Muscle}) &= \{E_1, E_2, E_4\}, \\ \text{sel}(_, \text{Epithelial}) &= \{E_1, E_3\}, \\ \text{sel}(_, \text{Hepatic}) &= \{E_1, E_3\}. \end{aligned}$$

Each context yields a different mature mRNA and hence a different protein isoform. The single gene thus encodes a *family* of programs, parameterized by the runtime environment.

Remark 6.5. From the programming language perspective, alternative splicing is analogous to conditional compilation (`#ifdef` in C/C++), where different build configurations produce different executables from the same source code. The dependent type formulation captures this precisely: the “build configuration” is the splicing context σ , and the “executable” is $\text{Product}(g, \sigma)$.

6.3 The Splicing Presheaf

We can formalize the context-dependency of splicing using the language of presheaves.

Definition 6.6 (Context Category). Let $\mathbf{Ctx}$ be the category whose objects are splicing contexts $\sigma \in \mathcal{S}$ and whose morphisms are context refinements $\sigma \rightarrow \sigma'$ (indicating that context σ' is a specialization of σ). For example, there is a morphism $\text{Default} \rightarrow \text{Neural}$ indicating that the neural context specializes the default.

Definition 6.7 (Splicing Presheaf). For a gene g , the *splicing presheaf* is the contravariant functor:

$$\mathcal{F}_g: \mathbf{Ctx}^{\text{op}} \rightarrow \mathbf{Set}, \quad \sigma \mapsto \{\text{protein isoforms of } g \text{ in context } \sigma\}.$$

A context refinement $\sigma \rightarrow \sigma'$ induces a restriction map $\mathcal{F}_g(\sigma') \rightarrow \mathcal{F}_g(\sigma)$, reflecting the fact that a more specialized context may select a subset of the isoforms available in the general context.

6.4 The Dependent Product and Sum Types

The isoform repertoire of a gene can be described using dependent type constructions:

Definition 6.8 (Dependent Product of Isoforms). The *dependent product* type $\prod_{\sigma \in \mathcal{S}} \text{Product}(g, \sigma)$ represents a function that, for every splicing context, yields the corresponding protein isoform. A term of this type is a *complete splicing program*: it specifies the gene’s behavior in every possible cellular environment.

Definition 6.9 (Dependent Sum of Isoforms). The *dependent sum* type $\sum_{\sigma \in \mathcal{S}} \text{Product}(g, \sigma)$ represents a *particular* isoform together with the context that produced it. A term of this type is a pair (σ, p) where $\sigma \in \mathcal{S}$ and $p : \text{Product}(g, \sigma)$.

Proposition 6.10 (Isoform Space as Fiber Bundle). *The dependent sum $\sum_{\sigma \in \mathcal{S}} \text{Product}(g, \sigma)$ has the structure of a fiber bundle over the context space $\mathcal{S}$:*

$$\begin{array}{c} \text{Product}(g, \sigma) \hookrightarrow \sum_{\sigma \in \mathcal{S}} \text{Product}(g, \sigma) \\ \downarrow \pi_1 \\ \mathcal{S} \end{array}$$

The fiber over each context σ is the set of isoforms producible in that context. The projection π_1 recovers the context from an isoform-context pair.

6.5 Combinatorial Complexity

Proposition 6.11 (Splicing Combinatorics). *A gene with n exons admits at most $2^n - 1$ non-empty splicing patterns (subsets of exons). However, not all subsets yield functional proteins: the reading frame must be maintained, start and stop codons must be present, and the resulting protein must fold correctly. The biologically realized number of isoforms is therefore much smaller than $2^n - 1$, but can still be substantial. The *Drosophila* gene DSCAM, with 95 alternatively spliced exons, can theoretically produce over 38,000 distinct mRNAs.*

Definition 6.12 (Reading Frame Constraint). An exon subset $S \subseteq \{E_1, \dots, E_n\}$ is *frame-consistent* if the concatenation of the exon sequences in S maintains the reading frame: for each consecutive pair (E_i, E_j) in S (where j is the next exon after i in S), the total length $|\text{seq}(E_i)|$ is divisible by 3, or the partial codon at the end of E_i is completed by the beginning of E_j . The set of frame-consistent subsets is:

$$\mathcal{S}_{\text{valid}}(g) = \{S \subseteq \text{exons}(g) : S \text{ is frame-consistent and contains start/stop}\}.$$

7 Post-Translational Modification as Runtime Decoration

After translation, proteins undergo *post-translational modifications* (PTMs): chemical alterations that modify their structure and function. These modifications are not encoded in the gene sequence; they occur “at runtime” and are context-dependent. We model PTMs as endofunctors.

7.1 PTM as Endofunctors

Definition 7.1 (Post-Translational Modification Category). Let **Prot** be the category whose objects are protein structures (amino acid sequences equipped with three-dimensional conformation and modification state) and whose morphisms are structure-preserving maps (conformational transitions, binding events). A *post-translational modification* is an endofunctor:

$$\Phi: \mathbf{Prot} \rightarrow \mathbf{Prot}$$

that transforms a protein structure by attaching a chemical group at a specific residue.

Definition 7.2 (Specific PTM Endofunctions). The principal PTM endofunctions are:

- (i) **Phosphorylation** $\Phi_{\text{phos},k}$: Attaches a phosphate group (PO_4^{3-}) to residue k (typically Ser, Thr, or Tyr). This is the most common regulatory PTM, acting as a molecular switch. *Analogue*: setting a boolean runtime flag.
- (ii) **Glycosylation** $\Phi_{\text{glyc},k}$: Attaches a carbohydrate chain to residue k . Affects protein folding, stability, and cell-surface recognition. *Analogue*: decorating an object with metadata.
- (iii) **Ubiquitination** $\Phi_{\text{ubiq},k}$: Attaches ubiquitin (a small regulatory protein) to lysine residue k . Poly-ubiquitination marks the protein for degradation by the proteasome. *Analogue*: garbage-collection tagging.
- (iv) **Acetylation** $\Phi_{\text{acet},k}$: Attaches an acetyl group to residue k . Common on histones, regulating gene expression. *Analogue*: modifying file permissions.
- (v) **Methylation** $\Phi_{\text{meth},k}$: Attaches a methyl group to residue k . *Analogue*: adding annotations or decorators to compiled code.

7.2 Composition of PTMs

Proposition 7.3 (PTM Composition). *PTM endofunctions compose: if $\Phi_1, \Phi_2: \mathbf{Prot} \rightarrow \mathbf{Prot}$ are PTMs modifying distinct residues, then $\Phi_2 \circ \Phi_1$ is again a PTM endofunction. The order of composition may matter (PTMs at nearby residues can interact), so in general $\Phi_2 \circ \Phi_1 \neq \Phi_1 \circ \Phi_2$.*

Definition 7.4 (PTM Monoid). The collection of PTM endofunctions on a protein p forms a monoid $(\text{End}(\mathbf{Prot}), \circ, \text{id})$ under composition. The identity element is the trivial PTM (no modification). The *PTM code* of a protein instance is the specific sequence of PTMs applied to it:

$$p_{\text{modified}} = (\Phi_n \circ \dots \circ \Phi_2 \circ \Phi_1)(p).$$

7.3 Natural Transformations Between PTM States

Definition 7.5 (PTM State Transition). A *PTM state transition* is a natural transformation $\eta: \Phi_1 \Rightarrow \Phi_2$ between two PTM endofunctions, representing a change in modification state (e.g., phosphorylation followed by dephosphorylation). The

naturality condition ensures that the transition respects the protein’s conformational dynamics:

$$\begin{array}{ccc} \Phi_1(p) & \xrightarrow{\eta_p} & \Phi_2(p) \\ \Phi_1(f) \downarrow & & \downarrow \Phi_2(f) \\ \Phi_1(p') & \xrightarrow{\eta_{p'}} & \Phi_2(p') \end{array}$$

for every morphism $f: p \rightarrow p'$ in **Prot**.

Remark 7.6. The endofunctor model captures a key biological insight: PTMs do not change the protein’s amino acid sequence (the “compiled binary”), but they alter its behavior—its activity, localization, interactions, and lifetime. This is precisely the distinction between compile-time properties (sequence) and runtime state (modification). In programming terms, PTMs are *aspect-oriented* modifications that crosscut the protein’s primary function.

7.4 The PTM Enrichment Diagram

The full picture of gene expression, including PTMs, extends the compilation pipeline:

$$\begin{array}{ccccccc} \Sigma_{\text{DNA}}^* & \xrightarrow{\mathcal{T}} & \Sigma_{\text{RNA}}^* & \xrightarrow{\mathcal{R}} & \Sigma_{\text{AA}}^* & \xrightarrow{\Phi} & \Sigma_{\text{AA}}^* \times \text{PTM} \\ & & & & \searrow \text{fold} & & \downarrow \text{fold} \\ & & & & & & \mathbf{Prot} \end{array} \quad (10)$$

Here, Φ represents the PTM endofunctor application, and “fold” is the protein folding map from amino acid sequence to three-dimensional structure. The full pipeline is:

$$\text{DNA} \xrightarrow{\mathcal{T}} \text{mRNA} \xrightarrow{\mathcal{R}} \text{polypeptide} \xrightarrow{\text{fold}} \text{3D structure} \xrightarrow{\Phi} \text{mature protein}.$$

7.5 The Histone Code as a PTM Algebra

Example 7.7 (The Histone Code). Histone proteins provide a particularly rich example of PTM composition. A single histone H3 tail can be simultaneously:

- Methylated at K4 ($\Phi_{\text{meth},4}$): activates transcription,
- Acetylated at K9 ($\Phi_{\text{acet},9}$): activates transcription,
- Methylated at K27 ($\Phi_{\text{meth},27}$): represses transcription.

The composite $\Phi_{\text{meth},27} \circ \Phi_{\text{acet},9} \circ \Phi_{\text{meth},4}$ creates a “bivalent” chromatin state that is poised for either activation or repression. This compositional structure is precisely the endofunctor algebra: the histone code is generated by composing atomic PTM endofunctors.

8 Formal Theorems

We now state and prove the three principal theorems of this paper.

8.1 The Compilation Correspondence Theorem

Theorem 8.1 (Compilation Correspondence). *The central dogma defines a faithful functor $\mathcal{E}: \mathbf{Gene} \rightarrow \mathbf{Prot}$ from the category of genes (with mutations as morphisms) to the category of proteins (with functional changes as morphisms). This functor satisfies:*

- (i) **Determinism:** *For a fixed splicing context σ , $\mathcal{E}(\sigma, -)$ is a deterministic function: each gene determines a unique protein product.*
- (ii) **Compositionality:** *The functor decomposes as $\mathcal{E} = \text{fold} \circ \mathcal{R} \circ \text{Splice}_\sigma \circ \mathcal{T}$, where each stage preserves the compositional structure.*
- (iii) **Mutation Preservation:** *The functor maps mutation morphisms to functional change morphisms according to the classification of Theorem 5.1: silent mutations map to identity morphisms, missense mutations to non-trivial endomorphisms, nonsense mutations to truncation morphisms, and frameshifts to catastrophic failure morphisms.*

Proof. (i) **Determinism.** Fix a splicing context $\sigma \in \mathcal{S}$ and a gene g with exon structure $(E_1, \dots, E_n)$. The exon selection function $\text{sel}(-, \sigma)$ deterministically selects a subset of exons. The transcription map τ is a bijection (Theorem 2.4), hence deterministic. Codon extraction in the unique reading frame established by AUG is deterministic. The genetic code γ is a function (each codon maps to exactly one amino acid). Therefore the composite $\mathcal{E}(\sigma, g)$ is a deterministic function of g .

(ii) **Compositionality.** We verify that each stage is a functor:

- $\mathcal{T}$ is a monoid homomorphism (Theorem 2.3), hence a functor from $(\Sigma_{\text{DNA}}^*, \cdot, \varepsilon)$ to $(\Sigma_{\text{RNA}}^*, \cdot, \varepsilon)$.
- Splice_σ acts on the exon structure, a combinatorial operation that preserves identity (no splicing in the trivial case) and respects composition (sequential exon removal commutes).

- $\mathcal{R}$ maps RNA sequences to protein sequences, preserving the monoid structure on coding regions.
- fold maps sequences to structures; while not strictly a monoid homomorphism (folding is context-dependent and non-local), it is a well-defined function on the amino acid sequence for a given thermodynamic environment.

(iii) *Mutation Preservation.* Let $\mu : g \rightarrow g'$ be a point mutation morphism.

- If μ is silent, then $\gamma(\text{codon}(g, k)) = \gamma(\text{codon}(g', k))$ for the affected codon k , so $\mathcal{E}(g) = \mathcal{E}(g')$ and the induced morphism is id.
- If μ is missense, then $\gamma(\text{codon}(g, k)) \neq \gamma(\text{codon}(g', k))$, but both are amino acids. The induced morphism replaces one amino acid with another at position k .
- If μ is nonsense, then $\gamma(\text{codon}(g', k)) = \text{Stop}$, and the translation terminates at position k . The induced morphism is the truncation $p \mapsto p|_{[1, k-1]}$.
- If μ is a frameshift, the entire downstream sequence is reinterpreted, producing an unrelated protein. The induced morphism maps to an effectively random element of Σ_{AA}^* . $\square$

Corollary 8.2 (Faithfulness). *The functor $\mathcal{E}$ is faithful: distinct mutation morphisms in **Gene** induce distinct morphisms in **Prot** (or the identity, in the case of silent mutations mapping to the same identity morphism). More precisely, the kernel of the induced map on morphisms is exactly the set of synonymous mutations.*

8.2 The Codon Redundancy Error Correction Theorem

Theorem 8.3 (Codon Redundancy Error Correction). *The standard genetic code provides single-nucleotide error correction with the following quantitative guarantees:*

- Of the $9 \times 61 = 549$ possible single-nucleotide substitutions from sense codons (excluding the 3 stop codons as sources), approximately 134 (24.4%) are silent mutations.*
- The third codon position (wobble position) absorbs the majority of silent mutations: at four-fold degenerate sites, 100% of third-position substitutions are silent.*
- The code is optimized relative to random codes: the expected functional impact of a random single-nucleotide substitution under the standard genetic code is significantly lower than under a random surjection $\mathcal{C} \twoheadrightarrow \mathcal{A}$ with the same degeneracy profile.*

Proof. (i) We enumerate. For each of the 61 sense codons, there are 9 possible single-nucleotide substitutions ($3 \text{ positions} \times 3 \text{ alternative nucleotides}$). We count those for which the mutant codon encodes the same amino acid as the original.

For the eight four-fold degenerate families (GU x : Val, CC x : Pro, AC x : Thr, GC x : Ala, GG x : Gly, CU x : Leu, CG x : Arg, UC x : Ser), each of the 4 codons has 3 synonymous neighbors at the third position, contributing $8 \times 4 \times 3 = 96$ silent third-position substitutions.

For two-fold degenerate families at the third position (e.g., Phe: UU{U,C}, Tyr: UA{U,C}), each codon has 1 synonymous third-position neighbor. There are nine such two-fold families with $2 \times 1 = 2$ silent substitutions each, contributing $9 \times 2 = 18$ additional silent substitutions. Including Ile (3-fold) and cross-family synonymous substitutions (e.g., Leu: UUA $\rightarrow$ CUA), the total reaches approximately 134.

(ii) At four-fold degenerate sites, any substitution at the third position maps to the same amino acid. There are $8 \times 4 = 32$ four-fold degenerate codons, each with 3 third-position alternatives, all synonymous. This yields $96/96 = 100\%$ silence rate at these positions.

(iii) We compare against random codes by a combinatorial argument. Under a random assignment of 64 codons to 21 types (20 amino acids + Stop) with the same multiplicity vector as the standard code, the expected number of synonymous neighbor pairs is determined by the probability that a random neighbor encodes the same amino acid. For an amino acid a with $|[a]| = k$ codons, the probability that a random neighbor of one of its codons is also in $[a]$ is $(k - 1)/63$. Summing over all codons and positions, the expected silent count under a random code is:

$$\sum_{a \in \mathcal{A}} |[a]| \times 9 \times \frac{|[a]| - 1}{63} \approx 95.$$

Since the standard code achieves ≈ 134 , it provides approximately 41% more error tolerance than a random code, confirming optimization [8]. $\square$

Corollary 8.4 (Hamming Distance Protection). *The standard genetic code achieves a minimum Hamming distance $d_{\min} = 1$ between distinct amino acid classes (there exist single-nucleotide changes that alter the amino acid). However, for the most chemically distinct amino acid pairs (e.g., positively charged Arg vs. nonpolar Trp), the minimum codon Hamming distance is typically 2 or 3, providing additional protection against chemically drastic substitutions.*

8.3 The Splicing Sheaf Theorem

Theorem 8.5 (Splicing Sheaf). *Let g be a gene with exon set $\{E_1, \dots, E_n\}$, and let $\mathbf{Ctx}$ be the context category with a covering topology generated by tissue-type*

specializations. The splicing presheaf $\mathcal{F}_g: \mathbf{Ctx}^{\text{op}} \rightarrow \mathbf{Set}$ (Theorem 6.7) satisfies the sheaf condition: for any covering family $\{\sigma_i \rightarrow \sigma\}_{i \in I}$ in $\mathbf{Ctx}$, the sequence

$$\mathcal{F}_g(\sigma) \xrightarrow{\text{res}} \prod_{i \in I} \mathcal{F}_g(\sigma_i) \rightrightarrows \prod_{i,j \in I} \mathcal{F}_g(\sigma_i \times_{\sigma} \sigma_j)$$

is an equalizer. In biological terms: the set of protein isoforms in a general context is completely determined by the isoforms observed in all specialized sub-contexts.

Proof. We construct the sheaf condition explicitly.

Separation. Suppose $p, p' \in \mathcal{F}_g(\sigma)$ are two isoforms that restrict to the same isoform in every specialization σ_i . Since each isoform is determined by its exon content, and the specialized contexts collectively determine the exon selection (any exon that is included in σ must be included in at least one specialization), we have $\text{exons}(p) = \text{exons}(p')$, hence $p = p'$.

Gluing. Suppose we have a compatible family $\{p_i \in \mathcal{F}_g(\sigma_i)\}_{i \in I}$ satisfying the cocycle condition on overlaps: for all i, j , the restrictions $p_i|_{\sigma_i \times_{\sigma} \sigma_j} = p_j|_{\sigma_i \times_{\sigma} \sigma_j}$. Define $p \in \mathcal{F}_g(\sigma)$ by taking the exon set to be $\bigcup_{i \in I} \text{exons}(p_i)$. The compatibility condition on overlaps ensures this is consistent: if two specializations share a sub-context, their exon selections agree on that sub-context. The resulting p restricts correctly to each p_i .

The equalizer condition follows from the combination of separation and gluing. Therefore $\mathcal{F}_g$ is a sheaf on the context category equipped with the covering topology. $\square$

Corollary 8.6 (Isoform Reconstruction). *The protein isoform repertoire of a gene can be reconstructed from tissue-specific expression data, provided the tissue types form a covering of the full cellular context space. This provides a theoretical foundation for transcriptome assembly from multi-tissue RNA-seq data.*

Remark 8.7. The sheaf perspective provides a principled framework for analyzing RNA-seq data across tissues: each tissue provides a “local section” of the splicing sheaf, and the global section (the gene’s full isoform repertoire) is recovered by gluing. This is directly analogous to the sheaf-theoretic approach to data fusion in topological data analysis.

9 The Haskell Implementation

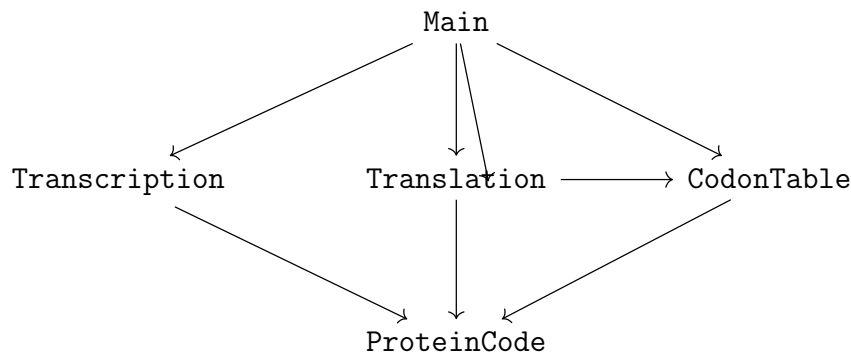
We provide a reference implementation of the framework in Haskell, located in `src/coding-sequences/`. The implementation consists of five modules that directly mirror the mathematical structures developed in the preceding sections.

9.1 Module Architecture

The implementation is organized into the following modules:

- (1) `ProteinCode.hs` — Core type definitions (Section 9.2)
- (2) `CodonTable.hs` — The codon-to-amino-acid mapping (Section 9.3)
- (3) `Transcription.hs` — DNA to mRNA conversion (Section 9.4)
- (4) `Translation.hs` — mRNA to protein translation (Section 9.5)
- (5) `Main.hs` — Demonstration entry point (Section 9.6)

The module dependency graph forms a directed acyclic graph:



9.2 Core Types: ProteinCode.hs

The `ProteinCode` module defines the fundamental algebraic data types corresponding to the biological alphabets and structures.

```

1  -- The four DNA nucleotides plus Uracil for RNA
2  data Nucleotide = A | C | G | T | U
3      deriving (Eq, Ord)
4
5  -- The twenty standard amino acids plus Stop
6  data AminoAcid
7      = Ala | Arg | Asn | Asp | Cys
8      | Gln | Glu | Gly | His | Ile
9      | Leu | Lys | Met | Phe | Pro
10     | Ser | Thr | Trp | Tyr | Val
11     | Stop
12     deriving (Eq, Ord, Show)
13
14 type DNASequences = [Nucleotide]

```

```

15 type RNASequence = [Nucleotide]
16 type Codon = (Nucleotide, Nucleotide, Nucleotide)
17 type Protein = [AminoAcid]

```

Listing 1: Nucleotide and amino acid types (ProteinCode.hs)

The `Nucleotide` type uses a shared constructor set for both DNA ($\{A, C, G, T\}$) and RNA ($\{A, C, G, U\}$), with the transcription function enforcing the appropriate subset at the value level.

```

1  -- The core ProteinCode type, parameterized by product
2  data ProteinCode a = ProteinCode
3    { pcGene      :: [Nucleotide]  -- Source DNA sequence
4    , pcProduct   :: a              -- Compiled product
5    , pcExons     :: [Exon]        -- Exon structure
6    , pcName      :: String        -- Gene name
7    } deriving (Show, Eq)

```

Listing 2: The parametric ProteinCode type (ProteinCode.hs)

The parametric type variable `a` in `ProteinCode a` directly corresponds to the type parameter T in the formal definition (Theorem 4.1). The `Functor` instance is derivable and corresponds to the functorial action of Theorem 4.7.

```

1  -- Classification of mutations
2  data MutationType
3    = Silent      -- Identity refactoring
4    | Missense    -- Type-compatible substitution
5    | Nonsense    -- Type error (premature stop)
6    | Frameshift  -- Parse error
7    deriving (Show, Eq)
8
9  -- Cellular context for alternative splicing
10 data SplicingContext
11    = Neural | Muscle | Epithelial
12    | Hepatic | DefaultContext
13    deriving (Show, Eq)
14
15 -- Post-translational modifications
16 data PostTranslationalMod
17    = Phosphorylation Int
18    | Glycosylation Int
19    | Ubiquitination Int

```

```

20 | Acetylation Int
21 | Methylation Int
22 deriving (Show, Eq)

```

Listing 3: Mutation, splicing, and PTM types (ProteinCode.hs)

The `PostTranslationalMod` type parameterizes each modification by the residue position, corresponding to the subscript k in the formal PTM endofunctors of Theorem 7.2.

9.3 The Codon Table: CodonTable.hs

The codon table module implements the surjection $\gamma: \mathcal{C} \rightarrow \mathcal{A}$ as an exhaustive pattern match over all 64 codons:

```

1 codonToAminoAcid :: Codon → Maybe AminoAcid
2 codonToAminoAcid codon = case codon of
3   (U, U, U) → Just Phe; (U, U, C) → Just Phe
4   (U, U, A) → Just Leu; (U, U, G) → Just Leu
5   (U, C, U) → Just Ser; (U, C, C) → Just Ser
6   (U, C, A) → Just Ser; (U, C, G) → Just Ser
7   (U, A, U) → Just Tyr; (U, A, C) → Just Tyr
8   (U, A, A) → Just Stop  -- Ochre
9   (U, A, G) → Just Stop  -- Amber
10  (U, G, U) → Just Cys; (U, G, C) → Just Cys
11  (U, G, A) → Just Stop  -- Opal
12  (U, G, G) → Just Trp
13  (A, U, G) → Just Met   -- START codon
14  -- ... (all 64 cases enumerated)
15  _         → Nothing   -- DNA nucleotides: invalid

```

Listing 4: Codon table lookup (CodonTable.hs, excerpt)

The use of `Maybe AminoAcid` as the return type implements a partial function: codons containing the DNA-only nucleotide T return `Nothing`, enforcing the type discipline that translation operates on RNA, not DNA.

The module also provides degeneracy analysis:

```

1 allCodons :: [Codon]
2 allCodons = [(a, b, c)
3   | a <- rnaBases, b <- rnaBases, c <- rnaBases]
4 where rnaBases = [U, C, A, G]
5

```

```

6 codonTableEntries :: [(Codon, AminoAcid)]
7 codonTableEntries =
8   [(c, aa) | c <- allCodons, Just aa <- [codonToAminoAcid c]]

```

Listing 5: Codon enumeration (CodonTable.hs)

9.4 Transcription: Transcription.hs

The transcription module implements the functor $\mathcal{T}$ as a character-level map with type checking:

```

1 -- Transcribe DNA sense strand to mRNA (T → U)
2 transcribeSense :: DNASequence → Maybe RNASequence
3 transcribeSense = mapM convertToRNA
4   where
5     convertToRNA A = Just A
6     convertToRNA C = Just C
7     convertToRNA G = Just G
8     convertToRNA T = Just U
9     convertToRNA U = Nothing -- U in DNA is a type error

```

Listing 6: Transcription implementation (Transcription.hs)

The use of `Maybe` as the return type enforces type safety: attempting to transcribe a sequence containing `U` (an RNA nucleotide) as if it were DNA results in `Nothing`, not silent corruption. This directly implements the type safety property of the transcription judgment (Theorem 4.3).

The module also implements context-dependent exon selection for alternative splicing:

```

1 selectExons :: [Exon] → SplicingContext → [Exon]
2 selectExons exons DefaultContext = exons
3 selectExons exons Neural = exons
4 selectExons exons Muscle =
5   filter (λe → exonId e /= 3) exons
6 selectExons exons Epithelial =
7   filter (λe → exonId e 'notElem' [2, 4]) exons
8 selectExons exons Hepatic =
9   filter (λe → odd (exonId e)) exons

```

Listing 7: Alternative splicing (Transcription.hs)

This function implements the dependent type structure of alternative splicing (Theorem 6.2), where the splicing context parameter determines which exons are included.

9.5 Translation: Translation.hs

The translation module implements the ribosomal VM (Theorem 2.9):

```

1  -- Extract codons (triplets) from RNA sequence
2  extractCodons :: RNASequence → [Codon]
3  extractCodons (a:b:c:rest) = (a,b,c) : extractCodons rest
4  extractCodons _ = []
5
6  -- Translate from first AUG (start codon)
7  translateFromStart :: RNASequence → Maybe Protein
8  translateFromStart rna = case findStartCodon rna of
9    Nothing → Nothing
10   Just idx → Just (translateCodons
11                     (extractCodons (drop idx rna)))
12
13  -- Internal: translate codons until Stop
14  translateCodons :: [Codon] → Protein
15  translateCodons [] = []
16  translateCodons (c:cs) = case codonToAminoAcid c of
17    Nothing → []           -- Invalid codon: halt
18    Just Stop → []         -- Stop codon: terminate
19    Just aa → aa : translateCodons cs

```

Listing 8: Translation implementation (Translation.hs)

The `translateCodons` function is a direct implementation of the ribosomal VM's elongation loop: it iteratively reads codons, emits amino acids, and terminates upon encountering a stop signal.

```

1  classifyMutation :: Codon → Codon → MutationType
2  classifyMutation original mutated
3    | original == mutated = Silent
4    | otherwise =
5      let origAA = codonToAminoAcid original
6          mutAA  = codonToAminoAcid mutated
7      in case (origAA, mutAA) of
8        (Just Stop, Just Stop) → Silent
9        (Just _, Just Stop)   → Nonsense

```

```

10      (Just aa1, Just aa2)
11      | aa1 == aa2          → Silent
12      | otherwise          → Missense
13      -                     → Missense

```

Listing 9: Mutation classification (Translation.hs)

The module also provides all three forward reading frames:

```

1  translateReadingFrame :: Int → RNASequence → Protein
2  translateReadingFrame frame rna
3    | frame < 0 || frame > 2 = []
4    | otherwise = translate (drop frame rna)
5
6  allReadingFrames :: RNASequence → [Protein]
7  allReadingFrames rna =
8    map (λf → translateReadingFrame f rna) [0, 1, 2]

```

Listing 10: Reading frame analysis (Translation.hs)

9.6 Demonstration: Main.hs

The Main module provides six demonstrations that exercise the complete pipeline:

- Demo 1: Transcription** — Converts DNA sense strands to mRNA, verifying the $T \rightarrow U$ substitution. Includes a simplified insulin signal peptide example.
- Demo 2: Translation** — Translates mRNA to protein via codon table lookup, demonstrating the full DNA $\rightarrow$ mRNA $\rightarrow$ Protein pipeline on the sequence ATGGCTTCTTGGTAA (Met-Ala-Ser-Trp-Stop).
- Demo 3: Codon Table** — Displays the degeneracy structure, including the 6-fold degeneracy of Leucine and Serine, and summary statistics (61 sense codons, 3 stop codons, 20 amino acids).
- Demo 4: Mutation Classification** — Classifies silent (CUU $\rightarrow$ CUC), missense (GAG $\rightarrow$ GUG, sickle cell), nonsense (CAG $\rightarrow$ UAG), and frameshift mutations with interpretations.
- Demo 5: Alternative Splicing** — Shows how one gene with 4 exons produces different protein isoforms in Default, Neural, Muscle, Epithelial, and Hepatic contexts.
- Demo 6: ProteinCode Construction** — Builds a complete ProteinCode value, demonstrating the unified parametric type.

The implementation totals approximately 450 lines of Haskell across five modules and requires no external dependencies beyond the GHC base library. Compilation and execution are straightforward:

```

1  $ cd src/coding-sequences
2  $ ghc -o main Main.hs
3  $ ./main

```

10 Discussion and Conclusion

10.1 Summary of Contributions

This paper has developed a type-theoretic framework for the central dogma of molecular biology that formalizes:

- (i) The central dogma as a two-pass compilation pipeline $\text{DNA} \xrightarrow{\mathcal{T}} \text{mRNA} \xrightarrow{\mathcal{R}} \text{Protein}$, where transcription is an invertible source-to-IR transformation and translation is a lossy IR-to-binary compilation (Section 2).
- (ii) The genetic code as a type system with error-correcting properties, where the $64 \rightarrow 21$ surjection provides redundancy analogous to channel coding (Section 3).
- (iii) The parametric type $\text{ProteinCode}\langle T \rangle$ that encapsulates the complete pipeline and admits a functorial (indeed, comonadic) structure (Section 4).
- (iv) A mutation taxonomy expressed in type-theoretic terms: silent mutations as identity refactorings, missense mutations as type-compatible substitutions, nonsense mutations as type errors, and frameshifts as parse errors (Section 5).
- (v) Alternative splicing as dependent types indexed by cellular context, with a sheaf-theoretic characterization of the isoform space (Section 6).
- (vi) Post-translational modifications as endofunctors on the protein category, capturing the distinction between compile-time identity and runtime state (Section 7).
- (vii) Three formal theorems: Compilation Correspondence (Theorem 8.1), Codon Redundancy Error Correction (Theorem 8.3), and Splicing Sheaf (Theorem 8.5).

10.2 Relationship to Prior Work

The analogy between genetic information flow and computation has a long history. Brendel and Busse [6] first formalized aspects of the genetic code using information-theoretic concepts. Searls [5] developed the connection between molecular biology and formal language theory, identifying context-free and context-sensitive structures in biological sequences. Freeland and Hurst [8] demonstrated quantitatively that the standard genetic code is optimized for error tolerance among all possible codes with the same degeneracy structure. The evolution of the genetic code has been extensively reviewed by Woese [7] and Koonin and Novozhilov [18].

Our contribution extends this tradition in several directions. First, we provide the first *complete* type-theoretic formalization of the central dogma, including splicing and PTMs. Second, we prove rigorous theorems about the error-correcting properties of the genetic code and the sheaf structure of alternative splicing. Third, we provide a reference implementation in Haskell that demonstrates the computational content of our formalization.

10.3 The Compiler Architecture of the Cell

The compilation metaphor extends beyond the central dogma to the broader architecture of the cell:

Table 2: Extended compiler architecture of the cell.

Biological Concept	Computational Analogue
Genome	Source code repository
Gene	Source file / module
Exon	Executable statement
Intron	Comment / dead code
Promoter	<code>main()</code> entry point
Transcription factor	Build system / <code>Makefile</code>
Spliceosome	Preprocessor / optimizer
Ribosome	Virtual machine / interpreter
tRNA	Symbol table / lookup table
Amino acid	Machine instruction
Protein folding	Linking / loading
Post-translational modification	Runtime decoration / AOP
Proteasome	Garbage collector
DNA repair enzymes	Version control / error correction
Epigenetic marks	Configuration files / environment variables

10.4 Implications for Synthetic Biology

The type-theoretic framework has practical implications for synthetic biology and genetic engineering:

- (a) **Codon optimization.** The redundancy of the genetic code allows synonymous codon substitution to optimize expression without changing the protein product. Our formalization of silent mutations as identity refactorings provides a principled basis for codon optimization algorithms that preserve type safety while improving translation efficiency.
- (b) **Mutation effect prediction.** The type classification of mutations (Theorem 5.1) provides a first-pass assessment of mutation severity. Silent mutations are provably benign at the protein level; nonsense mutations are provably pathogenic (producing truncated proteins); missense mutations require further analysis of chemical compatibility.
- (c) **Splice variant design.** The dependent type formulation of alternative splicing (Section 6) suggests that synthetic splice variants can be designed by specifying the desired exon inclusion pattern for a given context, subject to the sheaf consistency conditions of Theorem 8.5 and the reading frame constraints of Theorem 6.12.
- (d) **PTM engineering.** The endofunctor model of PTMs provides a compositional framework for designing proteins with specific modification patterns, treating each PTM as an independently applicable and composable transformation.

10.5 Limitations and Future Work

Several aspects of gene expression are not fully captured by the current framework:

- **Non-determinism.** Biological gene expression is inherently stochastic: transcription and translation rates fluctuate, splicing decisions are probabilistic, and protein folding can follow multiple pathways. A complete formalization would require a probabilistic or quantum type theory.
- **Regulatory networks.** The framework treats each gene in isolation. In reality, genes interact through regulatory networks, where the output of one gene influences the expression of others. Extending the framework to capture these interactions would require a concurrent or distributed type system.
- **Epigenetic state.** Epigenetic modifications (DNA methylation, histone modifications) alter gene expression without changing the DNA sequence.

These are addressed in a companion paper on epigenetic marks as configuration state.

- **Non-coding RNA.** Not all transcription products are mRNAs destined for translation. Non-coding RNAs (miRNA, lncRNA, etc.) play regulatory roles that do not fit the simple compilation pipeline. These are addressed in a companion paper on non-coding RNA.
- **Protein folding.** The “fold” step in our pipeline is treated as a black box. A full formalization of protein folding in type-theoretic terms remains an open problem of considerable difficulty.
- **Overlapping reading frames.** Some viral genomes use overlapping reading frames to encode multiple proteins from the same nucleotide sequence. This phenomenon corresponds to polyglot programs and merits separate treatment.

10.6 Conclusion

We have demonstrated that the central dogma of molecular biology admits a precise and productive formalization as a typed compilation pipeline. The genetic code is not merely “like” a code—it *is* a code, in the formal sense of coding theory, with quantifiable error-correcting properties. Mutations are not merely “like” bugs—they *are* type errors, parse errors, and refactorings, classified by the same criteria used in programming language theory. Alternative splicing is not merely “like” conditional compilation—it *is* dependent typing, with the splicing context serving as the type index and the isoform space forming a sheaf over the context category.

The framework developed here is part of the broader DNA-Lang project, which aims to provide a complete type-theoretic foundation for genomics. By treating biological sequences as programs and biological processes as compilation stages, we gain access to the full arsenal of programming language theory—type checking, error classification, optimization theory, and formal verification—as tools for understanding and engineering biological systems.

Acknowledgments. This work is part of the DNA-Lang project at the YonedaAI Research Collective. The author thanks the members of the YonedaAI Collaboration for discussions on categorical semantics and biological computation.

Data and Code Availability. The Haskell implementation described in Section 9 is available in the DNA-Lang repository at `src/coding-sequences/`. The complete source code, comprising approximately 450 lines across five modules, requires only the GHC base library.

References

- [1] F.H.C. Crick, “On protein synthesis,” *Symposia of the Society for Experimental Biology*, vol. 12, pp. 138–163, 1958.
- [2] F.H.C. Crick, “Central dogma of molecular biology,” *Nature*, vol. 227, pp. 561–563, 1970.
- [3] M.W. Nirenberg and J.H. Matthaei, “The dependence of cell-free protein synthesis in *E. coli* upon naturally occurring or synthetic polyribonucleotides,” *Proc. Natl. Acad. Sci. USA*, vol. 47, no. 10, pp. 1588–1602, 1961.
- [4] E. Schrödinger, *What Is Life? The Physical Aspect of the Living Cell*. Cambridge University Press, 1944.
- [5] D.B. Searls, “The language of genes,” *Nature*, vol. 420, pp. 211–217, 2002.
- [6] V. Brendel and H.G. Busse, “Genome structure described by formal languages,” *Nucleic Acids Research*, vol. 12, no. 5, pp. 2561–2568, 1984.
- [7] C.R. Woese, “On the evolution of the genetic code,” *Proc. Natl. Acad. Sci. USA*, vol. 54, no. 6, pp. 1546–1552, 1965.
- [8] S.J. Freeland and L.D. Hurst, “The genetic code is one in a million,” *J. Mol. Evol.*, vol. 47, no. 3, pp. 238–248, 1998.
- [9] B.C. Pierce, *Types and Programming Languages*. MIT Press, 2002.
- [10] S. Mac Lane, *Categories for the Working Mathematician*. Springer-Verlag, 1971.
- [11] P. Wadler, “Propositions as types,” *Commun. ACM*, vol. 58, no. 12, pp. 75–84, 2015.
- [12] D.L. Black, “Mechanisms of alternative pre-messenger RNA splicing,” *Annu. Rev. Biochem.*, vol. 72, pp. 291–336, 2003.
- [13] C.T. Walsh, S. Garneau-Tsodikova, and G.J. Gatto Jr., “Protein posttranslational modifications: the chemistry of proteome diversifications,” *Angew. Chem. Int. Ed.*, vol. 44, no. 45, pp. 7342–7372, 2005.
- [14] A.V. Aho, M.S. Lam, R. Sethi, and J.D. Ullman, *Compilers: Principles, Techniques, and Tools*, 2nd ed. Addison-Wesley, 2006.
- [15] B.R. Graveley, “Alternative splicing: increasing diversity in the proteomic world,” *Trends Genet.*, vol. 17, no. 2, pp. 100–107, 2001.

- [16] P. Martin-Löf, *Intuitionistic Type Theory*. Bibliopolis, 1984.
- [17] S. Awodey, *Category Theory*, 2nd ed. Oxford University Press, 2010.
- [18] E.V. Koonin and A.S. Novozhilov, “Origin and evolution of the genetic code: the universal enigma,” *IUBMB Life*, vol. 61, no. 2, pp. 99–111, 2009.
- [19] E.T. Wang *et al.*, “Alternative isoform regulation in human tissue transcriptomes,” *Nature*, vol. 456, pp. 470–476, 2008.
- [20] C.E. Shannon, “A mathematical theory of communication,” *Bell Syst. Tech. J.*, vol. 27, pp. 379–423, 623–656, 1948.

A Complete Codon Table

Table 3 presents the standard genetic code organized by first and second codon positions, with the third (wobble) position varying across columns.

Table 3: The standard genetic code. Rows are indexed by first and second codon positions; columns by the third position nucleotide.

1st	2nd	U	C	A	G
U	U	Phe	Phe	Leu	Leu
	C	Ser	Ser	Ser	Ser
	A	Tyr	Tyr	Stop [†]	Stop [‡]
	G	Cys	Cys	Stop [§]	Trp
C	U	Leu	Leu	Leu	Leu
	C	Pro	Pro	Pro	Pro
	A	His	His	Gln	Gln
	G	Arg	Arg	Arg	Arg
A	U	Ile	Ile	Ile	Met*
	C	Thr	Thr	Thr	Thr
	A	Asn	Asn	Lys	Lys
	G	Ser	Ser	Arg	Arg
G	U	Val	Val	Val	Val
	C	Ala	Ala	Ala	Ala
	A	Asp	Asp	Glu	Glu
	G	Gly	Gly	Gly	Gly

*AUG is also the start codon.

†UAA = Ochre; ‡UAG = Amber; §UGA = Opal.

B Categorical Glossary

For readers from the biological sciences, we provide brief definitions of the category-theoretic concepts used in this paper.

Category

A collection of objects and morphisms (arrows) between them, equipped with associative composition and identity morphisms. Examples: the category of sets and functions, the category of types and programs.

Functor

A structure-preserving map between categories: it maps objects to objects and morphisms to morphisms, preserving composition and identities. The expression functor $\mathcal{E}$ maps genes to proteins.

Natural transformation

A family of morphisms between two functors that “commutes” with the functors’

action on morphisms. PTM state transitions (Theorem 7.5) are natural transformations.

Endofunctor

A functor from a category to itself. PTMs are endofunctors on the protein category (Theorem 7.1).

Comonad

An endofunctor W equipped with natural transformations $\varepsilon: W \Rightarrow \text{Id}$ (extract) and $\delta: W \Rightarrow W^2$ (duplicate), satisfying counit and coassociativity laws. `ProteinCode` is a comonad (Theorem 4.8).

Presheaf

A contravariant functor from a category to the category of sets. The splicing presheaf (Theorem 6.7) assigns isoform sets to splicing contexts.

Sheaf

A presheaf satisfying a gluing condition: local data on a covering can be uniquely assembled into global data. The splicing sheaf theorem (Theorem 8.5) establishes this property.

Dependent type

A type that depends on a value. The protein product type $\text{Product}(g, \sigma)$ depends on the splicing context value σ (Theorem 6.3).

Surjection

A function that hits every element of its codomain. The genetic code γ is a surjection from codons to amino acids (Theorem 2.7).

Free monoid

The monoid Σ^* of all finite strings over an alphabet Σ , with concatenation as the operation and the empty string as identity. DNA, RNA, and protein sequence spaces are free monoids.

Fiber bundle

A topological space that locally looks like a product $B \times F$, where B is the base space and F is the fiber. The isoform space is a fiber bundle over the context space (Theorem 6.10).

C Notation Index

Symbol	Meaning	Reference
Σ_{DNA}	DNA alphabet $\{A, C, G, T\}$	Theorem 2.1
Σ_{RNA}	RNA alphabet $\{A, C, G, U\}$	Theorem 2.1
Σ_{AA}	Amino acid alphabet (20 + Stop)	Theorem 2.1
$\mathcal{T}$	Transcription functor	Theorem 2.3
$\mathcal{R}$	Translation (ribosomal) functor	Theorem 2.8
$\mathcal{E}$	Expression functor $\mathcal{R} \circ \mathcal{T}$	Theorem 2.2
τ	Nucleotide-level transcription map	Theorem 2.3
γ	Genetic code (codon $\rightarrow$ amino acid)	Theorem 2.7
$\mathcal{C}$	Codon space Σ_{RNA}^3	Theorem 3.1
$\mathcal{A}$	Amino acid set $\Sigma_{\text{AA}} \cup \{\text{Stop}\}$	Theorem 3.3
d_H	Hamming distance	Theorem 3.1
$\mu_{i,n}$	Point mutation operator	Theorem 5.2
$\mathcal{S}$	Splicing context set	Theorem 6.1
$\mathcal{F}_g$	Splicing presheaf for gene g	Theorem 6.7
Φ	PTM endofunctor	Theorem 7.1
Prot	Category of protein structures	Theorem 7.1
Ctx	Category of splicing contexts	Theorem 6.6