



Desarrollo de un prototipo de una plataforma web modular para la centralización, análisis y alerta temprana del bajo rendimiento en nadadores juveniles, caso de estudio: Fabio Toro Academy, Cali - Colombia

Brian León Hoyos
Joyd Esteban Lasprilla Collazos

Proyecto de grado entregado para obtener el título de
Magister en Ingeniería de Software

Dirigida por
MSc. Juan Carlos Martínez Arias

Pontificia Universidad Javeriana Cali
Facultad de Ingeniería y Ciencias
Maestría en Ingeniería de Software
Santiago de Cali
05 de junio de 2026

Santiago de Cali, 05 de junio de 2026.

Señores

Pontificia Universidad Javeriana Cali.

Ph.D. Luisa Rincón

Directora Maestría en Ingeniería de Software.

Cali.

Cordial Saludo.

Por medio de la presente hago constar que en mi calidad de director de trabajo de grado he revisado el proyecto titulado “Desarrollo de un prototipo de una plataforma web modular para la centralización, análisis y alerta temprana del bajo rendimiento en nadadores juveniles, caso de estudio: Fabio Toro Academy, Cali - Colombia” realizado por los estudiantes de Magister en Ingeniería de Software Brian León Hoyos (cod: 9030253) y Joyd Esteban Lasprilla Collazos (cod: 9012390), el cual se encuentra terminado y considero que cumple con los requisitos para ser sustentado.

Atentamente,

**Juan Carlos
Martínez
Arias**

Firmado digitalmente
por Juan Carlos
Martínez Arias
Fecha: 2026.06.04
16:24:26 -05'00'

MSc. Juan Carlos Martínez Arias

Santiago de Cali, 05 de junio de 2026.

Señores

Pontificia Universidad Javeriana Cali

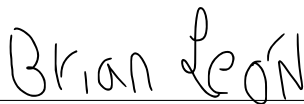
Ph.D. Luisa Rincón

Directora Maestría en Ingeniería de Software
Cali.

Cordial Saludo.

Nos permitimos presentar a su consideración el proyecto de grado titulado “Desarrollo de un prototipo de una plataforma web modular para la centralización, análisis y alerta temprana del bajo rendimiento en nadadores juveniles, caso de estudio: Fabio Toro Academy, Cali - Colombia” con el fin de cumplir con los requisitos exigidos por la Universidad y para que sea sometido a revisión del jurado y cumpla su aprobación, para conseguir posteriormente el título de Magister en Ingeniería de Software.

Atentamente,



Brian León Hoyos
Código: 9030253



Joyd Esteban Lasprilla Collazos
Código: 9012390

Ficha Resumen

Trabajo de Grado Maestría en Ingeniería de Software

TÍTULO: Desarrollo de un prototipo de una plataforma web modular para la centralización, análisis y alerta temprana del bajo rendimiento en nadadores juveniles. Caso de estudio: Fabio Toro Academy, Cali - Colombia

1. Énfasis: Ingeniería de Software
2. Área de trabajo: Ingeniería de Software / Tecnología Deportiva / Sistemas de Información / Análisis de Datos y Aprendizaje Automático.
3. Tipo de proyecto (Aplicado, Innovación, Investigación): Innovación
4. Estudiantes: Brian León Hoyos / Joyd Esteban Lasprilla Collazos.
5. Correos electrónicos: brianleon@javerianacali.edu.co / joyd@javerianacali.edu.co.
6. Dirección y teléfono: Calle 13B Bis # 48-29 - 3122944868 / Calle 8C # 46-61 - 3148066262.
7. Director: Juan Carlos Martínez Arias.
8. Vinculación del director: Profesor de planta, Facultad de Ingeniería y Ciencias.
9. Correo electrónico del director: juancmartinez@javerianacali.edu.co.
10. Co-Director (Si aplica):
11. Grupo o empresa que lo avala (Si aplica):
12. Otros grupos o empresas:
13. Palabras clave: desarrollo web, plataforma de análisis deportivo, detección de bajo rendimiento, natación juvenil, arquitectura modular, escalabilidad, modelo predictivo.
14. ODS que aplica al proyecto (Agenda 2030): ODS 3 (Salud y Bienestar), ODS 9 (Industria, Innovación e Infraestructura), ODS 10 (Reducción de las Desigualdades).
15. Fecha de inicio: 10 de enero del 2026
16. Resumen: Este proyecto de trabajo de grado de maestría abordó el desarrollo de un prototipo de plataforma web modular orientada a la centralización, el análisis, la predicción y la generación de alertas tempranas sobre el bajo rendimiento en nadadores juveniles. La problemática principal se relacionó con la gestión fragmentada y no estandarizada de los datos de entrenamiento, situación que dificultó el seguimiento del rendimiento, la planificación estratégica

y la comunicación efectiva entre entrenadores, deportistas y padres de familia, en el contexto del caso de estudio de Fabio Toro Academy, en Cali, Colombia.

La propuesta consistió en la materialización de un prototipo que incorporó una arquitectura de software modular, módulos de gestión y visualización de datos mediante tableros interactivos, y modelos analíticos y predictivos orientados a la detección temprana de patrones de bajo rendimiento. Los modelos se entrenaron en una primera fase con datos reales capturados en la academia y, ante su baja disponibilidad, en una segunda fase con un conjunto de datos sintético contextualizado; por ello sus resultados deben interpretarse como evidencia de viabilidad técnica y no como una caracterización empírica del rendimiento. En concreto, la clasificación del bajo rendimiento se abordó mediante tres algoritmos de clasificación, mientras que la predicción del tiempo de la siguiente sesión de entrenamiento se desarrolló a partir de un perceptrón multicapa. El sistema fue validado en Fabio Toro Academy como caso de estudio, evidenciando su alineación con las necesidades de los actores involucrados.

Como resultado, se desarrolló un prototipo funcional orientado a apoyar la toma de decisiones basadas en datos en el ámbito de la natación juvenil.

Agradecimientos

Ante todo, damos gracias a Dios por habernos dado la fortaleza, la sabiduría y la perseverancia necesarias para afrontar cada etapa de este proceso. Su guía fue el sostén más profundo en los momentos de dificultad y la luz que orientó cada decisión a lo largo del camino.

El desarrollo de este trabajo de grado no habría sido posible sin el apoyo, la dedicación y el compromiso de numerosas personas e instituciones que nos acompañaron en este recorrido. A todas ellas, nuestro más sincero y profundo agradecimiento.

A nuestras familias, por su paciencia incondicional, su respaldo emocional y por ser el motor que nos impulsó a continuar incluso en los momentos de mayor exigencia. Sus palabras de aliento y su confianza en nosotros fueron fundamentales para llegar a esta meta.

A nuestro director de trabajo de grado, MSc. Juan Carlos Martínez Arias, por su orientación permanente, su rigor académico y su disposición constante para guiar cada decisión técnica y metodológica del proyecto. Su acompañamiento fue determinante para la coherencia y calidad del resultado final.

A Fabio Toro Academy, en especial a su entrenador jefe, Fabio Toro Obando, atletas y padres de familia, por abrir las puertas de su institución y por su disposición genuina durante las sesiones de validación del prototipo. Su perspectiva práctica fue indispensable para que el sistema respondiera a necesidades reales.

Finalmente, este trabajo es también un reflejo de lo que aprendimos sobre nosotros mismos: que la perseverancia, la adaptación ante la incertidumbre y el trabajo en equipo son capacidades tan importantes como cualquier destreza técnica. Llevamos con nosotros no solo el conocimiento construido, sino la certeza de que la tecnología, cuando se diseña con propósito social, puede transformar realidades concretas.

Resumen

Este proyecto de trabajo de grado de maestría abordó el desarrollo de un prototipo de plataforma web modular orientada a la centralización, el análisis, la predicción y la generación de alertas tempranas sobre el bajo rendimiento en nadadores juveniles. La problemática principal se relacionó con la gestión fragmentada y no estandarizada de los datos de entrenamiento, situación que dificultó el seguimiento del rendimiento, la planificación estratégica y la comunicación efectiva entre entrenadores, deportistas y padres de familia, en el contexto del caso de estudio de Fabio Toro Academy, en Cali, Colombia.

La propuesta consistió en la materialización de un prototipo que incorporó una arquitectura de software modular, módulos de gestión y visualización de datos mediante tableros interactivos, y modelos analíticos y predictivos orientados a la detección temprana de patrones de bajo rendimiento. Los modelos se entrenaron en una primera fase con datos reales capturados en la academia y, ante su baja disponibilidad, en una segunda fase con un conjunto de datos sintético contextualizado; por ello sus resultados deben interpretarse como evidencia de viabilidad técnica y no como una caracterización empírica del rendimiento. En concreto, la clasificación del bajo rendimiento se abordó mediante tres algoritmos de clasificación, mientras que la predicción del tiempo de la siguiente sesión de entrenamiento se desarrolló a partir de un perceptrón multicapa. El sistema fue validado en Fabio Toro Academy como caso de estudio, evidenciando su alineación con las necesidades de los actores involucrados.

Como resultado, se desarrolló un prototipo funcional orientado a apoyar la toma de decisiones basadas en datos en el ámbito de la natación juvenil.

Palabras Clave: arquitectura modular, modelo analítico, plataforma web, prototipo, rendimiento deportivo, natación juvenil.

Abstract

This master's thesis project addressed the development of a modular web platform prototype aimed at centralizing, analyzing, predicting, and generating early warnings related to poor performance in youth swimmers. The main problem was associated with the fragmented and non-standardized management of training data, a situation that hindered performance monitoring, strategic planning, and effective communication among coaches, athletes, and parents. The project was conducted within the context of the Fabio Toro Academy case study in Cali, Colombia.

The proposed solution consisted of the implementation of a prototype that incorporated a modular software architecture, data management and visualization modules through interactive dashboards, and analytical and predictive models aimed at the early detection of poor performance patterns. In an initial phase, the models were trained using real data collected at the academy; however, due to the limited availability of such data, a contextualized synthetic dataset was used in a second phase. Therefore, the results should be interpreted as evidence of technical feasibility rather than as an empirical characterization of athletic performance. Specifically, poor performance classification was addressed using three classification algorithms, while the prediction of the time for the following training session was developed using a multilayer perceptron. The system was validated at Fabio Toro Academy as a case study, demonstrating its alignment with the needs of the stakeholders involved.

As a result, a functional prototype was developed to support data-driven decision-making in the field of youth swimming. **Keywords:** modular architecture, analytical model, web platform, prototype, athletic performance, youth swimming.

Índice general

Agradecimientos	7
1. Introducción	1
1.1. Definición del problema	2
1.1.1. Planteamiento del problema	2
1.1.2. Sistematización	3
1.2. Objetivos del proyecto	4
1.2.1. Objetivo General	4
1.2.2. Objetivos específicos	4
1.3. Delimitaciones y alcances	4
1.4. Justificación del trabajo de grado	5
1.5. Metodología de la investigación	6
1.5.1. Etapa 1: Identificación, recolección y procesamiento de datos relevantes	7
1.5.2. Etapa 2: Desarrollo del modelo analítico y predictivo	7
1.5.3. Etapa 3: Diseño y desarrollo del prototipo de la plataforma web modular	8
1.5.4. Etapa 4: Validación del prototipo en el caso de estudio	8
1.6. Resultados obtenidos	8
2. Marco de referencia	11
2.1. Marco teórico de referencia y antecedentes	11
2.1.1. Bases teóricas	11
2.1.2. Estado del Arte	17
3. Identificación, caracterización y preparación de datos de rendimiento en nadadores juveniles	21
3.1. Fase 1: recolección, preprocesamiento y caracterización de los datos reales	22
3.1.1. Recolección de datos	22
3.1.2. Justificación de las variables	23
3.1.3. Técnicas de preprocesamiento de datos	23
3.1.4. Caracterización del conjunto de datos reales	25
3.1.5. De la Fase 1 a la Fase 2: motivación de la transición	26
3.2. Fase 2: construcción y caracterización del conjunto sintético	27
3.2.1. Motivación metodológica	27
3.2.2. Supuestos y diseño de la generación	27
3.2.3. Estructura general del conjunto generado	28
3.2.4. Análisis exploratorio: verificación del generador	29

3.2.5. Alcance y limitaciones de la Fase 2	29
3.3. Trazabilidad de las variables y su rol en el modelado	29
3.4. Resumen del capítulo	30
4. Desarrollo del modelo analítico y predictivo para la detección temprana de bajo rendimiento deportivo	31
4.1. Fase 1: Construcción y evaluación del modelo predictivo de bajo rendimiento	31
4.1.1. Definición del parámetro de riesgo: bajo rendimiento	31
4.1.2. Modelos de clasificación seleccionados	32
4.1.3. Proceso de entrenamiento y selección de hiperparámetros	32
4.1.4. Resultados obtenidos y análisis	33
4.1.5. Conclusión y selección del mejor modelo	34
4.2. Fase 2: Predicción del rendimiento futuro mediante un modelo MLP	35
4.2.1. Motivación metodológica de la Fase 2	35
4.2.2. Preprocesamiento y selección de variables	35
4.2.3. Selección y arquitectura del modelo	35
4.2.4. Resultados del entrenamiento y evaluación	36
4.2.5. Alcance y limitaciones de la Fase 2	36
4.3. Resumen del capítulo	37
5. Diseño y desarrollo del prototipo de la plataforma web modular	39
5.1. Requerimientos del sistema	39
5.1.1. Diagrama de casos de uso	42
5.2. Arquitectura de la plataforma	44
5.2.1. Selección y justificación tecnológica	45
5.2.2. Frontend: React con TypeScript y Vite	46
5.2.3. Backend: Spring Boot con Java	46
5.2.4. Servicio de analítica: Python con FastAPI y scikit-learn	46
5.2.5. Base de datos: PostgreSQL	47
5.2.6. Síntesis de la decisión tecnológica	47
5.3. Fase 1	47
5.3.1. Enfoque arquitectónico.	47
5.3.2. Arquitectura en capas del backend.	48
5.3.3. Modelo de Representación Arquitectónica C4	48
5.3.4. Definición Arquitectónica del Frontend	51
5.3.5. Integración con el servicio de analítica y predicción	52
5.3.6. Modelo de Datos Relacional	53
5.3.7. Resultados de la Implementación de la Fase 1	54
5.3.8. Proyección hacia una Arquitectura Modular	56
5.4. Fase 2	57
5.4.1. Trazabilidad entre requerimientos y decisiones arquitectónicas	57

5.4.2. Enfoque arquitectónico	59
5.4.3. Arquitectura Monolito Modular basada en DDD	59
5.4.4. Integración del servicio de analítica y predicción	61
5.4.5. Modelo de Datos Relacional	62
5.4.6. Resultados de la Implementación de la Fase 2	63
5.4.7. Proyección a Escalamiento de la Arquitectura	67
5.4.8. Evaluación de la Evolución Arquitectónica	68
5.5. Resumen del capítulo	69
6. Validación del prototipo en el caso de estudio Fabio Toro Academy	71
6.1. Diseño de la validación	71
6.2. Validación de la Fase 1: prototipo inicial	72
6.2.1. Validación funcional con el entrenador jefe	72
6.2.2. Evaluación técnica del componente analítico	73
6.2.3. Síntesis de hallazgos y decisiones de rediseño	73
6.3. Validación de la Fase 2	74
6.3.1. Validación del conjunto de datos sintético contextualizado	74
6.3.2. Verificación del cumplimiento del objetivo analítico	74
6.3.3. Validación de la evolución arquitectónica	75
6.3.4. Evaluación del apoyo a la toma de decisiones y la gestión formativa	75
6.4. Verificación del cumplimiento de los objetivos	76
6.5. Limitaciones y amenazas a la validez	78
6.6. Resumen del capítulo	78
7. Conclusiones	79
7.1. Conclusiones	79
7.2. Trabajos futuros	81
7.3. Lecciones aprendidas	82
Bibliografía	85
A. Detalle de la recolección, el preprocesamiento y el conjunto de datos	89
A.1. Fase 1: detalle de la recolección y el preprocesamiento	89
A.1.1. Justificación del proceso de recolección	89
A.1.2. Diccionario completo de variables	89
A.1.3. Detalle de las etapas de preprocesamiento	90
A.1.4. Resultados del preprocesamiento sobre los datos reales	91
A.2. Fase 2: detalle de la construcción y el análisis del conjunto sintético	94
A.2.1. Supuestos de generación	94
A.2.2. Índice de bienestar: definición unificada	95
A.2.3. Flujo de construcción y validaciones automáticas	95

A.2.4. Análisis exploratorio detallado	96
B. Detalle técnico del modelado analítico y predictivo	101
B.1. Fase 1: detalle del modelo de clasificación de bajo rendimiento	101
B.1.1. Espacios de búsqueda de hiperparámetros	101
B.1.2. Definición formal de las métricas de evaluación	103
B.1.3. Matrices de confusión	104
B.1.4. Comparación gráfica de métricas	105
B.1.5. Curvas ROC	106
B.1.6. Curvas Precision-Recall	107
B.1.7. Análisis de importancia de variables	107
B.1.8. Análisis de sobreajuste	111
B.2. Fase 2: detalle del modelo MLP de predicción del rendimiento futuro	112
B.2.1. Pipeline de preprocesamiento del dataset sintético	112
B.2.2. Selección de variables predictoras	113
B.2.3. Comparación de alternativas y arquitectura del modelo	115
B.2.4. Interpretación gráfica de los resultados	117
C. Metodología de entrenamiento utilizada en Fabio Toro Academy	119
C.1. Propósito del anexo	119
C.2. Organización general del entrenamiento	119
C.3. Zonas y capacidades de entrenamiento	120
C.4. Programación del macrociclo diciembre-marzo	121
C.5. Relación de la metodología de entrenamiento con el prototipo desarrollado	123

Índice de figuras

3.1. Flujo metodológico de la gestión de datos a lo largo de las dos fases: recolección y caracterización de datos reales (Fase 1) y construcción y análisis de un conjunto sintético controlado (Fase 2).	22
3.2. Distribución de sesiones por nadador (etiquetas anonimizadas). La concentración en unos pocos atletas limita la capacidad de generalización inter-individual.	26
5.1. Diagrama caso de uso entrenador	43
5.2. Diagrama caso de uso acudiente	43
5.3. Diagrama caso de uso deportista	44
5.4. Diagrama arquitectura con stack tecnologico	45
5.5. Diagrama contexto	49
5.6. Diagrama contenedores	50
5.7. Diagrama componente backend - Fase 1	51
5.8. Diagrama componente frontend - Fase 1	52
5.9. Diagrama componente predicción - Fase 1	53
5.10. Modelo relacional de la plataforma en la Fase 1	54
5.11. Pantalla de autenticación de usuarios	55
5.12. Tablero de control para el entrenador	55
5.13. Formulario de registro de sesión de entrenamiento	56
5.14. Visualización de predicción, alertas y factores de riesgo	56
5.15. Diagrama de componentes del backend bajo arquitectura Monolito Modular basada en DDD	61
5.16. Modelo relacional de la plataforma en la Fase 2	63
5.17. Registro de sesión de entrenamiento con integración al módulo analítico.	65
5.18. Dashboard del entrenador con indicadores consolidados de rendimiento.	66
5.19. Módulo de alertas tempranas con las condiciones de riesgo identificadas por el modelo predictivo.	67
A.1. Distribución de registros por estilo (conjunto sintético).	97
A.2. Distribución de registros por tipo de sesión (conjunto sintético).	98
A.3. Distribución de la diferencia de tiempo (tiempo real menos planificado).	98
A.4. Relación entre el índice de bienestar y la diferencia de tiempo.	99
A.5. Evolución temporal de la diferencia media diaria.	99
A.6. Distribución de estado de ánimo y presencia de dolor (conjunto sintético).	100
B.1. Matrices de confusión para los tres modelos evaluados en el conjunto de prueba. . .	104

B.2. Comparación de las cinco métricas de clasificación para los tres modelos.	105
B.3. Curvas ROC de los tres modelos comparadas con el clasificador aleatorio.	106
B.4. Curvas Precision-Recall para los tres modelos.	107
B.5. Importancia de variables según Random Forest.	108
B.6. Coeficientes estandarizados de la Regresión Logística.	109
B.7. Importancia de variables según XGBoost (ganancia).	110
B.8. Análisis de <i>overfitting</i> : comparación de métricas en entrenamiento versus validación cruzada.	111
B.9. Curvas de entrenamiento y validación del modelo MLP para MAE y MSE.	117
B.10. Relación entre valores reales y predichos, y distribución de errores en el conjunto de prueba.	117
B.11. Error absoluto medio del modelo MLP por nadador y por estilo en el conjunto de prueba.	118

Índice de tablas

1.1. Resultados obtenidos	9
2.1. Comparativa de plataformas digitales para natación carreras juvenil.	18
3.1. Estructura del instrumento de recolección (Google Forms).	23
3.2. Pipeline de preprocesamiento.	24
3.3. Caracterización del conjunto de datos reales tras el preprocesamiento.	25
3.4. Relaciones inyectadas por diseño frente a propiedades emergentes en la generación sintética.	28
3.5. Indicadores estructurales del conjunto sintético (reproducibles con semilla 42).	28
3.6. Rol previsto de las variables en el modelado.	30
4.1. Criterios de riesgo para la construcción de la variable objetivo <i>bajo_rendimiento</i>	32
4.2. Métricas comparativas de los tres modelos en el conjunto de prueba.	33
4.3. Desempeño del modelo MLP en entrenamiento, validación y prueba.	36
5.1. Trazabilidad entre requerimientos y decisiones arquitectónicas	58
5.2. Comparación arquitectónica entre la Fase 1 y la Fase 2	68
6.1. Dimensiones del diseño de validación, evidencia e instrumentos.	72
6.2. Comparación cualitativa del proceso de seguimiento antes y después del prototipo.	76
6.3. Verificación del cumplimiento de los objetivos del proyecto.	77
A.1. Diccionario de variables del conjunto de datos.	90
A.2. Estadísticos descriptivos de las variables principales (datos reales, $n = 191$).	92
A.3. Sesiones por nadador (etiquetas anonimizadas, datos reales).	93
A.4. Distribución de variables categóricas (datos reales, $n = 191$).	94
A.5. Distribución por estilo y por tipo de sesión (sintético, asistentes, $n = 1559$).	96
A.6. Distribución de bienestar y dolor (sintético, asistentes, $n = 1559$).	97
B.1. Espacio de búsqueda de hiperparámetros para Regresión Logística.	102
B.2. Espacio de búsqueda de hiperparámetros para Random Forest.	102
B.3. Espacio de búsqueda de hiperparámetros para XGBoost.	103
B.4. Decisiones del pipeline de preprocesamiento de la Fase 2.	112
B.5. Columnas excluidas del modelado y motivo de exclusión.	113
B.6. Variables predictoras seleccionadas para el modelo MLP.	114
B.7. Comparación de alternativas evaluadas para el modelo de regresión.	115
B.8. Componentes principales de la arquitectura y entrenamiento del modelo MLP.	116

C.1. Componentes generales de la metodología de entrenamiento en Fabio Toro Academy	120
C.2. Nomenclatura de zonas y capacidades de entrenamiento	120
C.3. Macro ciclo utilizado como referencia en el caso de estudio	122
C.4. Distribución del volumen en agua por zona durante el macro ciclo	123

Introducción

La presente investigación aborda el desarrollo de un prototipo de plataforma web modular orientado a la centralización, análisis, predicción y generación de alertas tempranas sobre el bajo rendimiento en nadadores juveniles. Este tema surge de una problemática latente en la natación formativa en Colombia, donde la gestión de la información del desempeño deportivo se encuentra fragmentada y carece de estandarización. En el caso de estudio seleccionado, la academia Fabio Toro Academy, los registros de entrenamiento, pruebas físicas y evaluaciones técnicas se manejan actualmente en formatos dispersos, como archivos físicos o digitales no integrados, lo que dificulta el seguimiento continuo de los atletas, la planificación estratégica y la comunicación efectiva entre entrenadores, deportistas y padres de familia. La relevancia de este trabajo radica en la necesidad de superar estas barreras tecnológicas para optimizar el proceso formativo.

Para responder a este desafío, el objetivo general del trabajo de grado es desarrollar un prototipo de plataforma web que centralice estos datos y aplique un modelo analítico para detectar tempranamente caídas en el rendimiento, apoyando así la toma de decisiones. La investigación se guía por la pregunta central sobre cómo materializar una solución tecnológica de estas características en un entorno real. Para lograrlo, se han definido objetivos específicos que incluyen la identificación y caracterización de variables clave (cargas, fisiología, bienestar), el entrenamiento de modelos predictivos basados en inteligencia artificial, la construcción de una arquitectura de software modular y escalable, y finalmente, la validación del sistema en la academia Fabio Toro Academy para evaluar su impacto en la gestión del monitoreo deportivo.

La metodología adoptada para cumplir con estos propósitos es de carácter aplicado e innovador, estructurándose en cuatro fases secuenciales. Inicialmente, se realizó la identificación y limpieza de datos históricos y actuales; posteriormente, se desarrollo el modelo analítico utilizando técnicas de aprendizaje automático supervisado. Una tercera fase se enfoco en el diseño y construcción del prototipo web modular que integro estos modelos y tableros interactivos, culminando con una fase de validación funcional en el entorno real del caso de estudio. Este enfoque metodológico busco garantizar no solo la viabilidad técnica de la herramienta, sino también su utilidad práctica para los usuarios finales.

El impacto potencial de esta investigación es multidimensional, abarcando aspectos sociales, académicos y técnicos. Socialmente, busca proporcionar transparencia y confianza en el proceso formativo, mejorando la comunicación entre las familias y el cuerpo técnico. Desde una perspectiva académica, aporta evidencia novedosa sobre la aplicación de arquitecturas modulares y evidencia preliminar sobre la viabilidad de integrar modelos analíticos y predictivos interpretables en una plataforma web modular para el monitoreo del rendimiento en natación juvenil en Colombia.

Técnicamente, ofrece a los entrenadores herramientas para personalizar el entrenamiento y reducir riesgos de sobreentrenamiento, promoviendo trayectorias deportivas más sostenibles y abriendo oportunidades para la innovación tecnológica en el sector.

1.1. Definición del problema

En Colombia, la formación deportiva en disciplinas acuáticas como la natación carreras ha adquirido especial relevancia en los últimos años, impulsada por políticas estatales, inversión institucional y resultados destacados en escenarios internacionales [de la Hoz, 2025](#).

El Ministerio del Deporte ha consolidado una apuesta estructural para fortalecer la reserva deportiva nacional y, junto a la Federación Colombiana de Natación (FECNA), ha priorizado la identificación y el acompañamiento de jóvenes talentos mediante programas como Escuelas de Talentos. Esta iniciativa beneficia a miles de atletas y promueve estándares competitivos que, según los lineamientos oficiales, requieren procesos de seguimiento técnico, control y desarrollo sistemático del rendimiento juvenil [Mindeporte, 2025](#).

Sin embargo, estos avances coexisten con retos estructurales en la gestión de la información del desempeño deportivo, donde aún predomina la fragmentación en la recolección y el análisis de datos, y las tecnologías de monitoreo especializado son de limitada adopción fuera de los centros élite.

En Cali, Colombia, Fabio Toro Academy, institución dedicada a la formación de 14 nadadores de entre 13 y 21 años provenientes de Cali y Panamá, enfrenta importantes limitaciones en el seguimiento del rendimiento deportivo de sus atletas. La información clave relacionada con las cargas de entrenamiento, los tiempos en pruebas, la percepción del esfuerzo y el bienestar de los deportistas se encuentra dispersa en distintos formatos, tanto impresos como digitales, principalmente en archivos no integrados de Excel y Word.

Además, estas prácticas carecen de criterios estandarizados que permitan análisis integrales y sistematizados para apoyar la planificación, la toma de decisiones y la comunicación entre padres de familia, entrenadores y nadadores, a pesar de la existencia de directrices nacionales y marcas mínimas competitivas a nivel federativo [FECNA, 2024](#). Esta situación limita el potencial competitivo de los deportistas y reduce la confianza en el proceso formativo, una problemática coincidente con la realidad de numerosas academias deportivas independientes en Colombia y América Latina, donde la carencia de sistemas estructurados de seguimiento y comunicación dificulta la consolidación de trayectorias deportivas y el acompañamiento efectivo al joven atleta.

1.1.1. Planteamiento del problema

La gestión del rendimiento deportivo juvenil ha adquirido una relevancia creciente, hasta el punto de generar dispositivos de predicción del conteo de patadas en natación estilo libre, buscando obtener más datos para predecir el rendimiento deportivo [Bianchi et al., 2022](#). Los sistemas de monitoreo de atletas han demostrado ser esenciales para prevenir el síndrome de sobreentrenamiento y optimizar el rendimiento, especialmente en deportistas jóvenes, donde las tasas de prevalencia

pueden alcanzar entre el 29 % y el 37 % durante su carrera deportiva [Bernaciková et al., 2022](#), particularmente en disciplinas como la natación. La falta de seguimiento sistemático y de herramientas tecnológicas de apoyo dificulta la detección temprana de caídas en el rendimiento, afectando el desarrollo integral de los atletas y su permanencia en el deporte [Barry et al., 2024](#).

La transición de los métodos tradicionales como hojas de cálculo y registros manuales hacia soluciones digitales es una tendencia global que busca mejorar la toma de decisiones de entrenadores y atletas [Goudsmit et al., 2022](#).

La literatura reciente confirma que los factores de sobrecarga, el manejo inadecuado de la fatiga y la ausencia de alertas tempranas están vinculados a resultados por debajo del potencial y a la frustración de los jóvenes atletas [Carvalho et al., 2024](#). Asimismo, se ha documentado que la implementación de tableros de control y plataformas web facilita la interpretación de datos, mejora la coordinación entre entrenadores y familias, y reduce los conflictos derivados de expectativas no alineadas [Kaya and Senbel, 2023](#). Además, se ha identificado que una comunicación abierta entre padres y entrenadores es fundamental para evitar malentendidos que puedan dañar irreparablemente o incluso terminar la relación entre ambos [College of Education and Human Development, 2023](#).

Sin embargo, estas soluciones aún son incipientes en academias de formación en natación en contextos latinoamericanos, donde la madurez digital en sectores distintos al núcleo TIC sigue siendo baja y enfrenta diversas barreras de adopción y sostenibilidad, incluso en casos como Colombia [García et al., 2023](#).

Esta carencia de seguimiento sistemático justifica la necesidad de una solución tecnológica que aborde los desafíos de la detección temprana, la gestión de datos y la comunicación, con el fin de mitigar las frustraciones y tensiones en el proceso formativo.

1.1.2. Sistematización

1.1.2.1. Pregunta general

- ¿Cómo se puede materializar un prototipo de una plataforma web modular que centralice y analice los datos de rendimiento para la detección temprana del bajo rendimiento en nadadores juveniles y apoye la toma de decisiones y la comunicación entre entrenadores, deportistas y padres de familia de una academia de formación juvenil de natación?

1.1.2.2. Subpreguntas

- ¿De qué manera se debe procesar y presentar la información de rendimiento de forma clara, facilitando la interpretación y apropiación de los datos por parte de usuarios con diferentes niveles de conocimiento técnico y deportivo?
- ¿Cómo se puede procesar la información recopilada para la identificación y generación de alertas tempranas del bajo rendimiento deportivo?

- ¿Qué arquitectura de software se debe emplear para garantizar la integración de datos heterogéneos (entrenamiento, pruebas, bienestar) y permitir la incorporación futura de nuevas funcionalidades, asegurando la sostenibilidad técnica y la replicabilidad del sistema?
- ¿Cómo se puede mostrar el apoyo que brinda el prototipo en un entorno real, determinando su impacto en el monitoreo del rendimiento, la generación de alertas tempranas del proceso formativo en el caso de estudio: Fabio Toro Academy?

1.2. Objetivos del proyecto

1.2.1. Objetivo General

Desarrollar un prototipo de una plataforma web modular que permita la centralización de datos y la aplicación de un modelo analítico para la detección temprana del bajo rendimiento en nadadores juveniles, apoyando la toma de decisiones en el proceso formativo y la comunicación entre deportistas, entrenadores y padres de familia. Aplicándolo al caso de estudio: Fabio Toro Academy, Cali, Colombia.

1.2.2. Objetivos específicos

1. Identificar y caracterizar los datos relevantes relacionados con el rendimiento de nadadores juveniles (cargas de entrenamiento, indicadores fisiológicos, resultados de pruebas, bienestar subjetivo y registros de asistencia), garantizando su integridad, calidad y adecuación para el posterior desarrollo de modelos analíticos avanzados.
2. Desarrollar y entrenar modelos analíticos y predictivos, fundamentados en técnicas de inteligencia artificial y aprendizaje automático, capaces de detectar patrones de bajo rendimiento y predecir riesgos de estancamiento o sobreentrenamiento en función de los datos recolectados.
3. Desarrollar un prototipo de una plataforma web, fundamentado en una arquitectura modular, que integre el modelo analítico y permita la gestión, visualización y comunicación eficiente de los datos, incluyendo la notificación de alertas tempranas ante la detección de bajo rendimiento deportivo.
4. Validar la plataforma en el caso de estudio Fabio Toro Academy, evaluando su apoyo en el monitoreo del rendimiento deportivo, las alertas emitidas y la representación gráfica de los datos procesados por el modelo analítico y predictivo, con el fin de determinar su impacto en la toma de decisiones y la gestión formativa.

1.3. Delimitaciones y alcances

El alcance de este trabajo contempló el desarrollo de un prototipo de plataforma web, fundamentado en una arquitectura modular. Dicho prototipo integró los componentes esenciales: un módulo

de gestión y almacenamiento centralizado de los datos requeridos para la medición del rendimiento deportivo, un modelo analítico y predictivo orientado a la detección temprana del bajo rendimiento, un sistema de generación automática de alertas y tableros interactivos para la visualización clara de indicadores clave. Este prototipo se desarrolló dentro del período establecido para el proyecto de grado, garantizando un entregable alineado con los objetivos propuestos.

La validación en entorno real constituyó un eje central del alcance del proyecto. El prototipo fue validado en Fabio Toro Academy, ubicada en Cali, Colombia, mediante el uso de datos reales correspondientes a nadadores juveniles con edades entre los 13 y los 21 años. La interacción con entrenadores, deportistas y padres de familia permitió obtener retroalimentación directa sobre el potencial de la plataforma como apoyo en la toma de decisiones. Este proceso contribuyó al ajuste del modelo analítico y predictivo, garantizando que el prototipo respondiera de manera adecuada a los objetivos planteados.

El alcance del proyecto se delimitó al desarrollo y la evaluación de un prototipo dentro del entorno de la academia seleccionada, sin extenderse aún a una adopción masiva ni contemplar, en esta primera iteración, la integración con aplicaciones móviles nativas, plataformas externas de gestión deportiva o despliegues en infraestructura de producción a gran escala. De igual forma, no se incluyó el uso de hardware especializado, como sensores IoT o sistemas de visión por computador para la captura automatizada de datos. Estas exclusiones permitieron concentrar los esfuerzos en un núcleo funcional sólido y manejable en términos de tiempo y recursos, garantizando la factibilidad del proyecto. Entre las limitaciones técnicas consideradas se encontraron la calidad y el volumen de los datos disponibles, la curva de adopción tecnológica de los actores clave y la restricción temporal para realizar una validación más extensa.

Finalmente, este alcance estableció una base para futuras fases de expansión, en las que se contempló la integración con dispositivos externos de monitoreo y la interoperabilidad con plataformas deportivas de alcance regional e internacional. De igual manera, se proyectó la posibilidad de realizar despliegues en entornos productivos basados en la nube, con el propósito de escalar la solución hacia múltiples academias y, eventualmente, hacia otras disciplinas deportivas.

1.4. Justificación del trabajo de grado

La presente investigación se justificó por la necesidad de fortalecer los procesos de seguimiento y formación en la natación juvenil en Colombia, particularmente en academias independientes como Fabio Toro Academy, las cuales carecían de herramientas tecnológicas integrales para centralizar y analizar información crítica de los atletas. En el contexto identificado, la gestión fragmentada de los datos limitó la detección temprana de bajo rendimiento, generó frustración en los deportistas y sus padres de familia, y restringió la capacidad de los entrenadores para tomar decisiones informadas.

Desde una perspectiva social, la plataforma buscó proporcionar transparencia y confianza en el proceso formativo, al permitir que padres de familia, entrenadores y atletas accedieran a indicadores claros y objetivos, brindando apoyo en la toma de decisiones y fortaleciendo la comunicación entre los actores involucrados. Asimismo, la retroalimentación estructurada contribuyó a la motivación

de los nadadores al mostrar de manera tangible sus avances, aspecto que se relacionó con una mayor permanencia en la práctica deportiva juvenil [Goudsmit et al., 2022](#).

En el plano académico y científico, este proyecto generó evidencia sobre la aplicabilidad de arquitecturas modulares en sistemas deportivos, superando enfoques tradicionales y favoreciendo la sostenibilidad y la capacidad de evolución de la solución propuesta. De igual manera, aportó al campo del análisis deportivo al explorar el uso de modelos analíticos y predictivos explicables, validados en un entorno real de formación. La documentación técnica derivada del prototipo se constituyó en un insumo relevante para futuras investigaciones en ingeniería de software aplicada al deporte, especialmente en contextos latinoamericanos, donde aún se evidenciaba una implementación limitada de sistemas de esta naturaleza [Carvalho et al., 2024](#).

En el ámbito técnico y deportivo, la propuesta ofreció a los entrenadores una herramienta que facilitó la personalización de las sesiones, optimizó la planificación y redujo riesgos asociados al sobreentrenamiento, favoreciendo el desarrollo progresivo de los atletas [Barry et al., 2024](#). En este sentido, una monitorización eficaz contribuyó a reducir y mantener el riesgo de sobreentrenamiento por debajo del 5% [Gabbett, 2016](#). Esto incrementó la eficiencia en la preparación deportiva y favoreció trayectorias más sostenibles para los atletas juveniles.

Finalmente, desde una perspectiva económica y de innovación, el proyecto abrió la puerta a oportunidades de emprendimiento tecnológico en el sector deportivo, al validar la factibilidad de un sistema aplicable a múltiples academias y adaptable a otros deportes. De este modo, la solución no se limitó a atender una necesidad puntual, sino que ofreció bases sólidas para la transferencia de conocimiento y la innovación en ingeniería de software aplicada al deporte.

1.5. Metodología de la investigación

La metodología adoptada para este trabajo de grado se fundamentó en un enfoque aplicado e innovador, orientado al diseño y desarrollo de un prototipo de plataforma web modular, validado en el caso de estudio Fabio Toro Academy e integrado con un modelo analítico y predictivo sustentado en técnicas de inteligencia artificial.

Adicionalmente, el proceso metodológico se desarrolló bajo un enfoque iterativo por fases. En este sentido, las cuatro etapas descritas no se ejecutaron de manera única y secuencial, sino que conformaron un ciclo completo que se llevó a cabo en una primera fase (Fase 1), permitiendo obtener un prototipo inicial y resultados preliminares.

Posteriormente, a partir de los hallazgos, lecciones aprendidas y oportunidades de mejora identificadas en la Fase 1, se ejecutó una segunda iteración del ciclo metodológico (Fase 2), en la cual se volvieron a desarrollar las cuatro etapas, incorporando ajustes y optimizaciones orientadas a mejorar la calidad del modelo, la funcionalidad de la plataforma y la pertinencia de los resultados obtenidos. Este enfoque permitió fortalecer progresivamente la solución propuesta, asegurando un proceso de mejora continua basado en evidencia.

El proceso metodológico se estructuró en cuatro etapas, organizadas en ciclos iterativos de desarrollo, cada una alineada con los objetivos específicos del proyecto, lo que garantizó coherencia

entre las actividades desarrolladas y los resultados esperados.

1.5.1. Etapa 1: Identificación, recolección y procesamiento de datos relevantes

Esta primera fase tuvo como propósito identificar y caracterizar los datos clave relacionados con el rendimiento de nadadores juveniles. Las actividades desarrolladas incluyeron:

- La revisión documental y la realización de entrevistas con los actores involucrados para definir las variables relevantes, tales como cargas de entrenamiento, indicadores fisiológicos, resultados de pruebas, bienestar subjetivo y asistencia.
- La recolección de datos históricos y actuales provenientes de diversas fuentes.
- La limpieza, transformación y estandarización de los datos con el fin de garantizar su calidad, integridad y consistencia.
- La definición de indicadores clave de rendimiento y de criterios preliminares para la detección de bajo rendimiento.
- La validación de los datos y de los indicadores definidos con los actores involucrados.

Este proceso contempló además:

- La limpieza y estandarización de la información, la cual se encontraba dispersa en diferentes formatos, como archivos de Excel, documentos físicos y registros no centralizados.
- La definición de indicadores clave de rendimiento, que sirvieron como base para la construcción de los modelos analíticos.
- La validación de las variables consideradas con entrenadores y demás actores involucrados, garantizando su pertinencia dentro del contexto del estudio.

1.5.2. Etapa 2: Desarrollo del modelo analítico y predictivo

En esta fase se abordó la construcción de modelos orientados a detectar patrones de bajo rendimiento y predecir riesgos de estancamiento o sobreentrenamiento. Las actividades desarrolladas incluyeron:

- La selección y aplicación de técnicas de inteligencia artificial y aprendizaje automático.
- El entrenamiento y la validación de los modelos utilizando los datos recolectados en la fase anterior.
- La evaluación del desempeño de los modelos mediante métricas relevantes.

- La aplicación de técnicas de explicabilidad, con el propósito de garantizar la transparencia y la comprensión de las predicciones generadas.
- La documentación de los modelos desarrollados y de los resultados obtenidos.

1.5.3. Etapa 3: Diseño y desarrollo del prototipo de la plataforma web modular

Esta fase tuvo como finalidad la construcción de un prototipo funcional de la plataforma web. Las actividades contempladas fueron:

- El diseño de la arquitectura modular de la plataforma, permitiendo la integración y actualización de componentes sin afectar la estabilidad del sistema.
- El desarrollo de módulos para la gestión de datos, la visualización de indicadores, la integración del modelo analítico y la emisión de alertas automáticas.
- El desarrollo de tableros interactivos adaptados a los diferentes perfiles de usuario, incluyendo entrenadores, nadadores y padres de familia.

1.5.4. Etapa 4: Validación del prototipo en el caso de estudio

Esta fase tuvo como objetivo evaluar el desempeño del prototipo en el caso de estudio. Las actividades desarrolladas incluyeron:

- La implementación de la plataforma en Fabio Toro Academy.
- La evaluación funcional del sistema, mediante la verificación de la correcta operación de los módulos, la generación de alertas y la visualización de los datos.
- El análisis de los resultados obtenidos y la elaboración de un informe final, incluyendo recomendaciones orientadas a futuras mejoras y a la escalabilidad del sistema.

1.6. Resultados obtenidos

El desarrollo permitió obtener resultados en cinco dimensiones principales: gestión y preparación de datos, modelado analítico y predictivo, desarrollo del prototipo funcional, evolución arquitectónica y validación en el caso de estudio. En conjunto, estos resultados evidencian la viabilidad técnica, funcional y metodológica de una plataforma web modular orientada a la centralización, análisis y generación de alertas tempranas sobre el bajo rendimiento en nadadores juveniles de Fabio Toro Academy. No obstante, los resultados predictivos de la Fase 2 deben interpretarse como una demostración controlada del funcionamiento del pipeline analítico y de su integración con la plataforma, y no como una caracterización empírica definitiva del rendimiento real de los nadadores.

Tabla 1.1: Resultados obtenidos

	Resultado obtenido	Interpretación	Alcance o limitación
Datos y preparación	Se caracterizaron variables de carga externa, carga interna, rendimiento técnico, bienestar subjetivo y contexto deportivo. En la captura inicial de la Fase 1 se consolidaron 216 registros, en plazo de dos meses. Posteriormente, ante la cobertura desigual de registros y finalizar la fase 1, se construyó un dataset sintético contextualizado para la Fase 2, con 1.680 registros, 17 columnas, 14 nadadores y 1.559 sesiones con asistencia efectiva.	La Fase 1 permitió validar que el instrumento de captura podía recoger información relevante del proceso de entrenamiento, pero también evidenció que la adherencia sostenida al registro y la estandarización de los datos eran condiciones críticas para alimentar modelos predictivos. La Fase 2 permitió disponer de un conjunto controlado para probar de forma reproducible el flujo de preprocesamiento, entrenamiento, validación e integración del modelo.	El dataset sintético no reemplaza los datos reales ni constituye evidencia empírica definitiva del rendimiento deportivo. Su valor principal es metodológico y técnico.
Modelado analítico y predictivo	En la Fase 1 se entrenaron y compararon tres modelos supervisados para clasificar bajo rendimiento: Regresión Logística, Random Forest y XGBoost. XGBoost obtuvo el mejor equilibrio para el sistema de alertas, con accuracy de 0,8214, precisión de 0,7778, recall de 0,7000 y F1-Score de 0,7368. En la Fase 2 se incorporó un Perceptrón Multicapa (MLP) para estimar el tiempo real esperado del nadador en la sesión siguiente, alcanzando en prueba un MAE de 3,688 segundos, un RMSE de 4,547 segundos y un R^2 de 0,778.	Los resultados sugieren que la plataforma puede abordar el bajo rendimiento desde dos enfoques complementarios: una alerta binaria para identificar posibles casos de riesgo y una estimación continua del rendimiento esperado en sesiones futuras.	Las métricas del MLP corresponden a un escenario controlado con datos sintéticos. Por tanto, deben entenderse como evidencia de viabilidad técnica y no como precisión definitiva en condiciones reales.
Prototipo funcional	Se desarrolló un prototipo web compuesto por un frontend en React con TypeScript, un backend en Spring Boot con Java, una base de datos PostgreSQL y un servicio de analítica en Python con FastAPI. El sistema implementó autenticación por roles, registro de sesiones de entrenamiento, tableros diferenciados para entrenador, deportista y acudiente, generación automática de alertas y exportación de reportes en PDF.	El prototipo materializa el paso de registros dispersos hacia una plataforma centralizada, consultable y orientada al apoyo de la toma de decisiones. Además, integra el componente analítico dentro de un flujo funcional comprensible para los actores del caso de estudio, permitiendo visualizar alertas, factores de riesgo e indicadores de seguimiento.	El resultado corresponde a un prototipo funcional validado en un caso de estudio, no a una solución desplegada masivamente en producción.
Evolución arquitectónica	La arquitectura evolucionó desde un monolito por capas en la Fase 1 hacia un Monolito Modular basado en Domain-Driven Design (DDD) en la Fase 2. El backend fue organizado en módulos funcionales: usuarios y autenticación, sesiones de entrenamiento, analítica y predicción, dashboards y reportes, y configuración.	La reorganización por dominios funcionales permitió reducir el acoplamiento, aumentar la cohesión y localizar mejor los cambios dentro de cada módulo. Esta evolución fortalece atributos relevantes para una solución de ingeniería de software, como mantenibilidad, modificabilidad, escalabilidad funcional y capacidad de evolución.	La arquitectura queda preparada para crecimiento futuro, pero no implica todavía una migración a microservicios ni un despliegue productivo de alta escala.

Continúa en la siguiente página

Continuación de la Tabla 1.1

	Resultado obtenido	Interpretación	Alcance o limitación
Validación en el caso de estudio	El prototipo fue validado en Fabio Toro Academy. El entrenador jefe reconoció como aportes concretos el mecanismo de alertamiento automático, la centralización del historial de sesiones y la generación de reportes estructurados. La validación técnica del componente analítico confirmó la pertinencia del enfoque multimodelo y orientó la evolución hacia el MLP en la Fase 2.	La validación permitió evidenciar utilidad percibida para el seguimiento individual de los nadadores, la interpretación conjunta de variables que antes estaban dispersas y el fortalecimiento de la comunicación entre entrenador, deportistas y acudientes mediante dashboards y reportes.	No se midieron impactos deportivos longitudinales ni cambios de rendimiento durante varias temporadas. La evidencia corresponde a validación funcional, técnica y de utilidad percibida.

En síntesis, los resultados obtenidos muestran que el proyecto cumplió con el propósito de desarrollar un prototipo de plataforma web modular capaz de centralizar información de entrenamiento, integrar modelos analíticos y predictivos, generar alertas tempranas y apoyar la toma de decisiones en el contexto de Fabio Toro Academy. El principal aporte del trabajo no consiste en demostrar un impacto deportivo definitivo, sino en materializar una solución de ingeniería de software que articula gestión de datos, aprendizaje automático, arquitectura modular y validación funcional en un entorno real de formación deportiva. La principal limitación identificada fue la dependencia de una captura sostenida y estandarizada de datos reales por parte de los nadadores, por lo que la continuidad del sistema requiere fortalecer la adherencia al registro y acumular información empírica que permita recalibrar los modelos en futuras iteraciones.

Marco de referencia

2.1. Marco teórico de referencia y antecedentes

Un estudio realizado en centros de formación de atletas juveniles en Indonesia resalta que la implementación de la ciencia deportiva y el monitoreo sistemático son factores clave para optimizar el desarrollo deportivo, aunque existen retos asociados a la infraestructura y la estandarización de herramientas de evaluación Penggalih et al., 2025. Este tipo de evidencia permite identificar la oportunidad de abordar el problema en contextos como la natación formativa y en academias como el caso de estudio en Fabio Toro Academy, en Cali, Colombia.

A continuación, se presentan los principales conceptos y modelos teóricos que sustentan este trabajo.

2.1.1. Bases teóricas

2.1.1.1. Plataformas web en el deporte juvenil

Las plataformas web han transformado la gestión y el análisis de datos en el ámbito deportivo, permitiendo la centralización de información y la interacción eficiente entre entrenadores, atletas y padres. Estas soluciones tecnológicas facilitan la toma de decisiones informadas y la colaboración entre los diferentes actores involucrados en el proceso formativo Goudsmit et al., 2022. En el contexto de la natación juvenil, la adopción de plataformas web contribuye a superar la fragmentación de datos y mejora la trazabilidad del rendimiento deportivo.

La relevancia de estas plataformas radica en su capacidad para integrar módulos funcionales, como la recolección de datos, la visualización de resultados y la generación de alertas, en un entorno accesible y escalable. Así, se promueve un entorno colaborativo y transparente, donde la información fluye de manera estructurada y segura entre todos los participantes del proceso deportivo. No obstante, la adopción efectiva de estas tecnologías puede verse limitada por factores como la infraestructura tecnológica, la capacitación de los usuarios y la aceptación cultural, especialmente en contextos con menor madurez digital Kaya and Senbel, 2023.

Para que estas plataformas sean sostenibles y adaptables a largo plazo, es fundamental que su diseño se base en principios de modularidad y escalabilidad, permitiendo la integración de nuevas funcionalidades y la adaptación a las necesidades cambiantes de las organizaciones deportivas.

2.1.1.2. Software Modular y Escalable: Principios y Aplicaciones

El diseño de software modular y escalable es fundamental para el desarrollo de sistemas sostenibles en el tiempo. La modularidad implica dividir la aplicación en componentes independientes, cada uno responsable de una funcionalidad específica, lo que facilita el mantenimiento, la evolución y la integración de nuevas capacidades [Guntakandla, 2025](#). Por su parte, la escalabilidad permite que el sistema responda eficientemente al aumento de usuarios o volumen de datos, ya sea mediante la adición de recursos (escalabilidad horizontal) o la optimización de los existentes (vertical).

En el ámbito deportivo, la aplicación de estos principios permite que las plataformas evolucionen conforme cambian las necesidades de las organizaciones o se incorporan nuevas tecnologías, como módulos de analítica avanzada o integración con dispositivos externos. La literatura respalda que el diseño modular y la arquitectura escalable son estrategias efectivas para garantizar la adaptabilidad, la eficiencia operativa y la integración de múltiples fuentes de datos en sistemas de información deportiva [Cheng, 2022](#). De esta manera, se promueve la capacidad de respuesta ante futuros desafíos.

La transición hacia arquitecturas modulares y escalables representa una mejora significativa respecto a los sistemas monolíticos tradicionales, que suelen presentar dificultades de mantenimiento y limitaciones para el crecimiento funcional [Guntakandla, 2025](#). Esta evolución tecnológica es clave para el éxito de soluciones orientadas al monitoreo y análisis del rendimiento deportivo juvenil. Sin embargo, la migración a este tipo de arquitecturas puede implicar desafíos relacionados con la interoperabilidad, la capacitación del personal y la gestión del cambio organizacional [Benjamin, 2025](#).

Dentro de estos principios de modularidad existen distintas formas de organizar la solución según cómo se agrupan sus componentes. Un enfoque ampliamente adoptado es el del *Monolito Modular*, que mantiene la aplicación como una única unidad desplegable pero organizada internamente en módulos definidos por áreas funcionales o dominios; esta organización suele apoyarse en el *Domain-Driven Design* (DDD), una aproximación que estructura el software a partir de los conceptos propios del dominio del problema. De esta manera se busca conservar la simplicidad de una solución unificada y, al mismo tiempo, favorecer la mantenibilidad y la evolución por áreas funcionales [Guntakandla, 2025](#); [Cheng, 2022](#).

Una de las formas más efectivas de implementar modularidad en la práctica es mediante el uso de APIs RESTful, las cuales han demostrado ser exitosas en plataformas deportivas y de salud por su capacidad de integración y actualización continua [Benjamin, 2025](#).

2.1.1.3. APIs RESTful

Las APIs RESTful (Representational State Transfer) constituyen un estilo arquitectónico ampliamente adoptado para el diseño de servicios web que facilitan la comunicación entre componentes de software. Este enfoque se basa en el protocolo HTTP y se centra en la manipulación de recursos identificados de forma única mediante URIs (Uniform Resource Identifiers) [López and Maya, 2024](#). A continuación, se presentan sus principios fundamentales y su relevancia para el presente proyecto.

Definición y principios fundamentales.

Una API REST define un conjunto de restricciones arquitectónicas que orientan el diseño de interfaces de comunicación entre sistemas. Los principios fundamentales incluyen: la identificación de recursos mediante URIs, donde cada entidad del dominio (por ejemplo, un nadador, una sesión de entrenamiento o una alerta) se representa como un recurso accesible a través de una dirección única; el uso de los verbos estándar del protocolo HTTP (GET, POST, PUT, DELETE) para realizar operaciones sobre dichos recursos; y la comunicación sin estado (*stateless*), lo que implica que cada solicitud del cliente contiene toda la información necesaria para ser procesada, sin que el servidor almacene contexto entre peticiones López and Maya, 2024; Contreras, 2020. Estas restricciones favorecen la simplicidad, la escalabilidad y la interoperabilidad del sistema.

La representación de los recursos se realiza comúnmente en formato JSON, un estándar ligero que facilita la lectura, la transmisión y el procesamiento de datos tanto por parte de los desarrolladores como de las aplicaciones cliente López and Maya, 2024. Frente a alternativas más pesadas como SOAP, REST simplifica la integración entre componentes heterogéneos y reduce la complejidad de las interfaces López and Maya, 2024.

Contratos, versionamiento y códigos de estado.

Las APIs actúan como contratos entre los distintos módulos del sistema, lo que permite versionar, probar y mantener cada componente de forma independiente Gonzalez et al., 2024; López and Maya, 2024. Los códigos de estado HTTP (como 200 para operaciones exitosas, 201 para creación de recursos, 400 para errores del cliente o 404 para recursos no encontrados) proporcionan una semántica estándar que facilita el diagnóstico de errores y la interacción entre el frontend y el backend de la plataforma.

Adicionalmente, la propiedad de idempotencia en operaciones como GET, PUT y DELETE garantiza que ejecutar la misma solicitud múltiples veces produzca el mismo resultado, lo que refuerza la confiabilidad y la tolerancia a fallos en la comunicación entre componentes Contreras, 2020.

Relevancia para el proyecto.

En el contexto de la plataforma desarrollada para Fabio Toro Academy, las APIs RESTful permiten estructurar la comunicación entre los módulos funcionales del sistema: la recolección de datos de entrenamiento, el procesamiento analítico mediante modelos de machine learning, la generación de alertas y la visualización de resultados para entrenadores, atletas y padres. Este diseño basado en recursos y operaciones estándar facilita la modularidad, la mantenibilidad y la eventual incorporación de nuevas funcionalidades sin comprometer la estabilidad del sistema Chamari et al., 2025; Benjamin, 2025.

2.1.1.4. Rendimiento Deportivo y Entrenamiento por Zonas en Natación

El rendimiento deportivo en natación se evalúa mediante indicadores como tiempos, volúmenes de entrenamiento y pruebas físicas. El entrenamiento por zonas implica segmentar la intensidad

del ejercicio (por ejemplo, zonas aeróbicas y anaeróbicas) para optimizar la adaptación fisiológica y prevenir el sobreentrenamiento [Carvalho et al., 2024](#). El monitoreo sistemático de estos indicadores es clave para detectar caídas de rendimiento y ajustar los planes de entrenamiento.

La integración de estos datos en la plataforma permite identificar patrones de bajo rendimiento y personalizar la intervención, apoyando la toma de decisiones basada en evidencia tanto para entrenadores como para atletas y padres. La metodología concreta de planificación y las zonas de entrenamiento empleadas por la academia, que sirvieron de referencia contextual para la caracterización de datos y el diseño del componente analítico, se detallan en el Anexo C.

Para maximizar el valor de los datos recolectados, es fundamental incorporar modelos analíticos y predictivos para la detección de bajo rendimiento deportivo, que permitan anticipar riesgos y actuar preventivamente.

2.1.1.5. Modelos Analíticos y Predictivos para la Detección de Bajo Rendimiento Deportivo

El uso de modelos analíticos y predictivos en el ámbito deportivo se fundamenta en la capacidad de procesar datos históricos y actuales para anticipar patrones de bajo rendimiento, riesgo de lesiones o sobreentrenamiento. En el presente proyecto, estas técnicas se apoyan en algoritmos de machine learning supervisado aplicados a datos tabulares, los cuales han demostrado ofrecer valor agregado en comparación con los enfoques descriptivos tradicionales, al permitir no solo visualizar lo ocurrido, sino proyectar posibles escenarios futuros [Cordeiro et al., 2025](#). En el contexto deportivo, esto implica identificar a tiempo factores como disminución en la velocidad, incremento en la fatiga percibida o reducción en la frecuencia de asistencia a entrenamientos.

Un aspecto relevante es la interpretabilidad de los modelos, dado que los entrenadores y padres requieren entender por qué el sistema clasifica a un nadador como “en riesgo” y qué variables inciden más en esa predicción [Carvalho et al., 2024](#). De esta manera, la plataforma no solo funciona como un sistema de alerta temprana, sino como una herramienta de apoyo pedagógico y de toma de decisiones informadas.

A partir de estas consideraciones, se seleccionaron las técnicas específicas de clasificación supervisada que se describen a continuación.

2.1.1.6. Técnicas de Machine Learning Supervisado aplicadas al Análisis Deportivo

El aprendizaje supervisado consiste en entrenar un modelo con datos etiquetados; es decir, cada instancia de entrenamiento incluye características (*inputs*) y una variable objetivo conocida (*output*). Una vez entrenado, el modelo predice la clase o valor de nuevas instancias [Jiang et al., 2020](#). Entre las familias de algoritmos más utilizadas en el análisis deportivo con datos tabulares se encuentran los modelos lineales, los ensambles de árboles de decisión y las redes neuronales artificiales, cada uno con fortalezas y limitaciones que justifican su estudio comparativo.

Regresión logística (Logistic Regression).

La regresión logística es un modelo lineal que estima la probabilidad de pertenencia a una clase mediante la función sigmoide. A pesar de su simplicidad matemática, constituye una línea base robusta y ampliamente utilizada en clasificación binaria [Shetty, 2023](#). Su principal fortaleza radica en la interpretabilidad: los coeficientes del modelo permiten identificar directamente qué variables aumentan o reducen la probabilidad de la clase objetivo y en qué magnitud, lo cual resulta especialmente valioso en contextos donde las decisiones del sistema deben ser comprensibles para usuarios no técnicos.

Bosques aleatorios (Random Forest).

Random Forest es un método de ensamble que construye múltiples árboles de decisión de forma aleatoria, cada uno entrenado con una muestra *bootstrap* del conjunto de datos y considerando subconjuntos aleatorios de variables en cada división. Las predicciones finales se obtienen mediante votación mayoritaria [Analytics Vidhya, 2024](#). Esta técnica reduce la varianza y mitiga el sobreajuste característico de los árboles individuales, ofreciendo un buen balance entre precisión y robustez frente al ruido en los datos. Además, proporciona estimaciones naturales de la importancia de cada variable a través de la reducción promedio de impureza (índice de Gini), lo que facilita la identificación de los factores más relevantes para la predicción.

XGBoost (Extreme Gradient Boosting).

XGBoost es un algoritmo de *boosting* por gradiente que construye árboles de decisión de forma secuencial, donde cada nuevo árbol corrige los errores residuales de los anteriores. A diferencia de Random Forest, que entrena árboles de forma independiente, XGBoost emplea un proceso iterativo de optimización con regularización integrada en la función objetivo, lo que le confiere un control más fino del sobreajuste. Este algoritmo ha demostrado resultados destacados en problemas de clasificación con datos tabulares de tamaño moderado, y su mecanismo de importancia basado en ganancia (*gain*) permite identificar qué variables contribuyen más a la reducción del error de clasificación.

Multilayer Perceptron (MLP).

El Multilayer Perceptron (MLP) es una red neuronal artificial de tipo *feedforward* compuesta por una capa de entrada, una o más capas ocultas y una capa de salida. Las capas ocultas aplican funciones de activación no lineales, lo que permite al modelo aproximar patrones complejos en problemas de clasificación y regresión sobre datos tabulares [Luo et al., 2025](#). A diferencia de los modelos lineales, el MLP captura interacciones no lineales entre variables, propiedad relevante en el análisis deportivo, donde el rendimiento suele depender de la combinación de factores fisiológicos, perceptuales y contextuales que rara vez se relacionan de forma estrictamente lineal [He and Choi, 2025](#). Su aplicación requiere, sin embargo, mecanismos de control del sobreajuste, validación adecuada y regularización, particularmente cuando el conjunto de datos es limitado, así como técnicas que aporten interpretabilidad a sus predicciones [Zhang et al., 2024](#).

La disponibilidad de estos cuatro modelos cubre un espectro complementario: desde un modelo lineal interpretable (Regresión Logística), pasando por un ensamble de *bagging* robusto (Random Forest) y un ensamble de *boosting* de alto rendimiento (XGBoost), hasta una red neuronal capaz de aproximar relaciones no lineales (MLP). Esta diversidad facilita la comparación rigurosa y fundamentada entre familias de algoritmos para problemas de clasificación y regresión con datos tabulares en el ámbito deportivo [Chicco and Jurman, 2020](#); [Jiang et al., 2020](#).

2.1.1.7. Sistemas de Alertas Tempranas

Un sistema de alerta temprana (SAT) en el ámbito deportivo se define como una herramienta que utiliza modelos predictivos y análisis de datos para anticipar riesgos asociados al bajo rendimiento, la fatiga o la aparición de lesiones, permitiendo intervenciones preventivas antes de que los problemas se manifiesten plenamente [Qian and Lee, 2025](#); [López-Valenciano et al., 2018](#).

La literatura demuestra que, incluso en ausencia de hardware especializado, es posible anticipar escenarios de riesgo mediante el análisis de datos tabulares (como registros de entrenamiento, cuestionarios de bienestar, métricas de rendimiento) y aplicando técnicas de machine learning supervisado [López-Valenciano et al., 2018](#). Esta aproximación resulta especialmente relevante en contextos con recursos limitados, donde la accesibilidad y la simplicidad de los datos son factores determinantes para la implementación de SAT efectivos.

Además de la capacidad técnica para generar alertas, la plataforma debe facilitar la comunicación efectiva entre entrenadores, atletas y padres, asegurando que las notificaciones sean comprensibles y accionables, promoviendo una cultura de prevención y respuesta coordinada [Qian and Lee, 2025](#). No obstante, la adopción de estos sistemas implica desafíos relacionados con la interpretación de alertas, la gestión de falsas alarmas y la integración en la rutina diaria de los equipos deportivos.

2.1.1.8. Comunicación de Actores: Entrenador, Atleta y Padres

La comunicación efectiva entre entrenadores, atletas y padres es fundamental para el éxito deportivo juvenil. Las plataformas digitales estructuradas facilitan la transparencia, la retroalimentación y la alineación de expectativas, reduciendo conflictos y mejorando la motivación y permanencia de los deportistas [Kaya and Senbel, 2023](#).

La tecnología debe facilitar canales de comunicación diferenciados y accesibles para cada actor, promoviendo un entorno colaborativo y de apoyo. La evidencia muestra que la comunicación regular y honesta mejora la confianza mutua y la cohesión en torno al atleta joven.

Finalmente, la protección de la información sensible es un aspecto crítico que debe ser considerado en el diseño de la plataforma.

2.1.1.9. Seguridad y Manejo de Información Sensible

En Colombia, la protección de información sensible en plataformas digitales está regulada principalmente por la Ley 1581 de 2012, que establece los principios y procedimientos para el tratamiento de datos personales en entidades públicas y privadas. Esta normativa exige la implementación de

medidas como el cifrado, el control de acceso y la gestión de incidentes para garantizar la confidencialidad, integridad y disponibilidad de los datos.

Además, el Ministerio TIC ha definido lineamientos específicos sobre seguridad y privacidad de la información, orientados a fortalecer la protección de los datos en entornos digitales [MinTIC, 2025](#).

En el ámbito deportivo, estas disposiciones se complementan con políticas internas de entidades como el Ministerio del Deporte, que enfatizan la importancia del consentimiento informado y la protección de datos sensibles de los atletas.

El cumplimiento de estas normativas es fundamental para generar confianza entre los usuarios y asegurar la sostenibilidad de la plataforma en el tiempo.

2.1.2. Estado del Arte

En esta sección se revisan algunas de las principales plataformas digitales utilizadas actualmente en la natación carreras juvenil. El objetivo es identificar sus características, ventajas y limitaciones para entender el estado actual de la tecnología en este campo.

2.1.2.1. Plataformas Digitales en Natación Juvenil

1. **SwimCloud** Es una plataforma digital orientada a la gestión integral de competencias de natación carreras. Permite la administración de inscripciones, resultados en tiempo real, rankings y gestión de equipos, además de ofrecer compatibilidad con archivos estándar como Hy-Tek y SwimTopia. La versión Pro incluye herramientas para la alineación automática de mejores marcas y sugerencias de alineaciones, facilitando la operación de competencias y la reducción de errores administrativos. Sin embargo, la documentación pública no evidencia módulos nativos para el monitoreo de carga de entrenamiento, bienestar o sistemas de alerta temprana, ya que su enfoque principal es la gestión de datos y resultados de competencias [SwimCloud, 2025](#).
2. **AthleteMonitoring** Es una plataforma SaaS enfocada en el monitoreo de atletas juveniles en deportes como la natación carreras. Ofrece seguimiento de bienestar, carga de entrenamiento (sRPE, ACWR, *strain/monotony*), lesiones, progresión y reportes personalizados, además de mensajería segura y cumplimiento de normativas internacionales de privacidad. Sus paneles personalizables y alertas visuales permiten identificar riesgos en tiempo real. No obstante, la documentación pública no describe funciones específicas para la gestión de competencias de natación, inscripciones o cronometraje, ya que su enfoque está en la salud, la carga y la prevención de lesiones [AthleteMonitoring.com, 2024](#).
3. **SwimTrack** Es una herramienta basada en visión por computador para el análisis automatizado de videos de carreras de natación. Permite la detección de nadadores, el cálculo de frecuencia y longitud de brazada, trayectorias y velocidad, utilizando grabaciones estándar de competencias. Esta solución ha sido utilizada en retos internacionales de análisis de video,

facilitando la extracción de métricas técnicas sin hardware adicional en piscina. Sin embargo, SwimTrack no está orientado a la gestión administrativa de clubes, monitoreo de carga o comunicación con familias, ya que su uso principal es el análisis biomecánico y técnico en contextos de alto rendimiento e investigación [MultimediaEval, 2022](#).

Herramienta	Funcionalidades claves	Ventajas	Limitaciones	Relevancia para el proyecto
SwimCloud	Gestión de competencias, inscripciones, resultados en vivo, rankings, administración de equipos; compatibilidad con archivos estándar; herramientas Pro para alineación automática.	Automatiza entradas y resultados; visualización en vivo; compatibilidad con formatos estándar; reduce errores administrativos.	No incluye monitoreo de carga/bienestar ni alertas tempranas.	Alta para operación de competencias y gestión de datos de rendimiento.
AthleteMonitoring	Monitoreo de bienestar, carga (sRPE, ACWR), lesiones, progresión; paneles personalizables, alertas, reportes; mensajería segura.	Alertas visuales y personalización de métricas; informes en tiempo real; enfoque en prevención.	No gestiona competencias de natación ni cronometraje.	Alta para monitoreo de carga, bienestar y prevención.
SwimTrack	Análisis de video de carreras: detección de nadadores, frecuencia/longitud de brazada, trayectorias/velocidad; conjuntos de datos para investigación.	Extrae métricas técnicas desde video estándar, sin hardware adicional; útil para análisis biomecánico.	No gestiona clubes, competencias ni monitorea carga/bienestar.	Alta para análisis técnico y biomecánico en natación carreras.

Tabla 2.1: Comparativa de plataformas digitales para natación carreras juvenil.

2.1.2.2. Estudios Académicos sobre Aplicaciones, Sistemas o Plataformas para Detección Temprana o Monitoreo del Rendimiento o Condiciones Relacionadas en Deportes

1. **Desarrollo de un sistema de monitoreo de rendimiento en tiempo real en deportes de equipo** [Aydn et al., 2025](#) Desarrollaron un sistema de monitoreo de rendimiento en tiempo real para deportes de equipo, específicamente aplicado en voleibol, el cual mide parámetros fisiológicos como frecuencia cardíaca, aceleración y temperatura central de varios atletas simultáneamente, transmitiendo los datos de forma inalámbrica mediante el protocolo ESP-NOW y mostrándolos en una interfaz gráfica accesible para los entrenadores. El sistema fue probado con seis atletas, demostrando que es capaz de recopilar y transmitir datos sin interrupciones incluso en ambientes interiores, destacándose por ser de bajo costo, accesible y eficiente en la reducción de pérdidas de datos. Este trabajo evidencia que ya existen plataformas aplicadas con capacidad de monitoreo continuo, lo que valida el concepto de detección

temprana del bajo rendimiento en un entorno deportivo real y aporta un referente práctico al estado del arte en esta línea de investigación.

2. **Validación de una aplicación basada en IA para evaluación del rendimiento deportivo** Aydemir et al., 2025 Estudio sobre la validez y confiabilidad de DeepSport, una aplicación basada en inteligencia artificial para la evaluación del rendimiento deportivo mediante análisis de imágenes, comparada frente al sistema de referencia OptoJump en la medición de saltos verticales (CMJ y SJ) en jugadores de baloncesto de élite. Los resultados del análisis estadístico mostraron una alta confiabilidad ($ICC > 0.90$) y ausencia de sesgos sistemáticos, concluyendo que DeepSport puede considerarse una herramienta válida y confiable para el análisis del rendimiento en contextos deportivos reales. Este estudio constituye una evidencia concreta de cómo las aplicaciones digitales con componentes de IA están siendo validadas académicamente para apoyar el análisis de desempeño, aportando al estado del arte desde la perspectiva de detección y seguimiento temprano de indicadores clave de rendimiento.

Identificación, caracterización y preparación de datos de rendimiento en nadadores juveniles

La calidad de un sistema analítico depende, antes que de los algoritmos, de la calidad de los datos que lo alimentan: los conjuntos con inconsistencias o valores fuera de rango propagan sus deficiencias a todo el modelado posterior (van Buuren, 2018; Foster et al., 2017). En la natación juvenil, una gestión rigurosa de los datos de entrenamiento contribuye a optimizar el rendimiento y a reducir el riesgo de lesiones (Barry et al., 2024); no obstante, en diversas academias de América Latina persisten brechas de captura y estandarización que dificultan la toma de decisiones basada en evidencia (Penggali et al., 2025).

Este capítulo desarrolla la identificación, recolección y procesamiento de los datos relevantes para el rendimiento. En términos del ciclo de vida estándar de un proyecto de aprendizaje, esta etapa corresponde a las fases de *comprensión de los datos* y *preparación de los datos*, previas y habilitantes del modelado. Coherente con el carácter iterativo de la metodología, el trabajo con los datos se ejecutó en dos fases. La **Fase 1** estableció la recolección, el preprocesamiento y la caracterización de los datos reales capturados en la academia; la **Fase 2**, surgida tras evaluar el prototipo inicial y los hallazgos de la Fase 1, construyó y caracterizó un conjunto de datos sintético controlado. La Figura 3.1 resume este flujo y el Anexo A amplía cada componente.



Figura 3.1: Flujo metodológico de la gestión de datos a lo largo de las dos fases: recolección y caracterización de datos reales (Fase 1) y construcción y análisis de un conjunto sintético controlado (Fase 2).

3.1. Fase 1: recolección, preprocesamiento y caracterización de los datos reales

3.1.1. Recolección de datos

La recolección se implementó mediante un formulario digital en Google Forms, diseñado para que los atletas lo completaran al finalizar cada sesión, dentro de los primeros treinta minutos. Esta decisión prioriza la captura oportuna de medidas subjetivas como el RPE (Percepción del esfuerzo), cuya validez se degrada a medida que aumenta el retraso del reporte (Saw et al., 2016). El instrumento consta de diecisiete ítems organizados en cinco secciones, según se resume en la Tabla 3.1.

Tabla 3.1: Estructura del instrumento de recolección (Google Forms).

Sección	Variables capturadas	Escala / formato
Identificación	Fecha, nadador	Fecha; lista desplegable
Datos del entrenamiento	Volumen total, minutos en zona aeróbica y anaeróbica, tipo de sesión, tiempos de 100 m	Metros; minutos; categórica (R1, R2, R1-2, VO2, RL, TL); segundos
Estado fisiológico	RPE, horas de sueño	RPE (escala del 1 al 10); horas
Bienestar emocional	Estado de ánimo, estrés, motivación	Ánimo (3 niveles); estrés 1–5; motivación 1–10
Contexto deportivo	Dolor, días desde competencia, fase del macrociclo	Sí/No; días; fase (4 niveles)

La elección de este método responde a cuatro criterios complementarios, inmediatez del registro, estandarización, centralización y accesibilidad sin hardware especializado, que en conjunto procuran un flujo de datos oportuno y trazable (Foster et al., 2017; Barry et al., 2024; Garcia et al., 2023).

3.1.2. Justificación de las variables

La selección de variables adopta un enfoque multidisciplinario que integra carga externa, carga interna y condicionantes del entorno. Esta articulación permite comprender con mayor precisión el rendimiento y facilita la detección temprana de estados de fatiga o sobreentrenamiento que comprometan la adaptación deportiva (Kellmann et al., 2018). Las variables se agrupan en cuatro dimensiones:

- **Carga de entrenamiento (externa):** volumen total nadado, minutos en zona aeróbica y anaeróbica, y tipo de sesión (R1, R2, R1-2, VO2, RL, TL).
- **Rendimiento técnico:** tiempo de 100 m planificado y tiempo de 100 m real logrado en la sesión.
- **Carga interna y bienestar:** RPE (escala del 1 al 10), horas de sueño, estado de ánimo, estrés percibido, motivación y presencia de dolor.
- **Contextuales:** días desde la última competencia y fase del macrociclo.

3.1.3. Técnicas de preprocesamiento de datos

El preprocesamiento condiciona la calidad de todo análisis posterior, por lo que se diseñó un pipeline de siete etapas que prioriza la reproducibilidad y la trazabilidad, de modo que cada transformación pueda auditarse y replicarse.

Tabla 3.2: Pipeline de preprocesamiento.

Etapa	Acción	Propósito
1	Estandarización de nomenclatura	Renombrar columnas a snake_case para una manipulación consistente.
2	Normalización de formatos numéricos	Unificar separadores del volumen nadado mediante reglas de dominio.
3	Conversión de tiempos	Llevar todos los formatos de 100 m a segundos.
4	Codificación de categóricas	Codificación ordinal y binaria según la naturaleza de cada variable.
5	Validación de rangos y fechas	Acotar (marcar como nulos) los valores y fechas fuera de los límites plausibles.
6	Tratamiento de valores ausentes	Distinguir nulos legítimos de información faltante.
7	Creación de variables derivadas	Construir el sRPE, el ratio aeróbico y el índice de bienestar.

De la séptima etapa surgen tres indicadores compuestos que se emplean en el análisis. Para garantizar su reproducibilidad se enuncian de forma explícita:

$$\text{sRPE} = \text{RPE} \times (\text{min. aeróbicos} + \text{min. anaeróbicos}) \quad (3.1)$$

$$\text{ratio aeróbico} = \frac{\text{min. aeróbicos}}{\text{min. aeróbicos} + \text{min. anaeróbicos}} \quad (3.2)$$

$$\begin{aligned} \text{índice de bienestar} = & 0,20 \frac{\text{motivación} - 1}{9} + 0,20 \left(1 - \frac{\text{estrés} - 1}{4} \right) \\ & + 0,20 \frac{\text{ánimo} - 1}{2} + 0,40 (1 - \text{dolor}) \end{aligned} \quad (3.3)$$

El *session-RPE* (Ecuación 3.1) integra la intensidad percibida con el volumen de la sesión; el *ratio aeróbico* (Ecuación 3.2) estima la polarización del entrenamiento; y el *índice de bienestar* (Ecuación 3.3) combina ánimo, motivación, estrés y dolor, ponderando el dolor con mayor peso (40 %) por su valor como señal de alarma ante posibles lesiones (Saw et al., 2016; Bittencourt et al., 2016; Kellmann et al., 2018). El índice está *normalizado* en el intervalo $[0, 1]$, donde 0 representa el estado más comprometido y 1 el más favorable, mediante la normalización de cada componente a esa escala. Esta es la *única* definición del índice empleada en todo el trabajo: el preprocesamiento de los datos reales y el generador del conjunto sintético (Fase 2) aplican exactamente la misma fórmula, de modo que el umbral de riesgo «índice de bienestar $\leq 0,3$ » (Capítulo 4) representa el mismo nivel de bienestar en ambos conjuntos.

3.1.4. Caracterización del conjunto de datos reales

La aplicación del pipeline sobre las respuestas crudas permite caracterizar el corpus efectivamente disponible. Más que enumerar cifras, interesa interpretar qué implican para la viabilidad del modelado. La Tabla 3.3 resume los indicadores principales.

Tabla 3.3: Caracterización del conjunto de datos reales tras el preprocesamiento.

Indicador	Valor
Respuestas crudas registradas	216
Sesiones válidas (con asistencia)	191
Tasa de asistencia	88,4 %
Nadadores con registros	13
Periodo plausible de captura	dic. 2025 – mar. 2026
Sesiones con tiempo de 100 m planificado ausente	46,6 %
Sesiones con tiempo de 100 m real ausente	49,7 %
Sesiones con ambos tiempos disponibles	96
Volumen medio por sesión	4.896 m
sRPE medio	418 UA
Sesiones con dolor reportado	26,2 %

Tamaño y cobertura. De las 216 respuestas se obtuvieron 191 sesiones con asistencia efectiva (88,4 %), correspondientes a 13 de los 14 nadadores de la academia; uno de los atletas no diligenció el formulario, de modo que no generó registros utilizables. El periodo plausible de captura abarca de diciembre de 2025 a marzo de 2026.

Disponibilidad del desenlace cronométrico. El hallazgo más determinante es que el tiempo de 100 m, variable central que el sistema busca anticipar, está ausente en cerca de la mitad de las sesiones (46,6 % el planificado y 49,7 % el real); solo 96 sesiones disponen de ambos tiempos válidos. La validación de rangos, ahora aplicada, descarta automáticamente los valores imposibles para una prueba de 100 m (por ejemplo, diferencias de varios cientos de segundos por errores de digitación), de modo que el desenlace restante es confiable. Sobre esas 96 sesiones, la diferencia entre el tiempo real y el planificado tiene una mediana de -2 s, indicando que el real suele ser ligeramente mejor que el planificado, y un 28,1 % de los casos corresponde a subrendimiento cronométrico (tiempo real superior al planificado). En la práctica, el desenlace que sostendría un modelo predictivo está disponible y es confiable en alrededor de un centenar de observaciones, base estrecha para entrenar un modelo cronométrico robusto.

Concentración por nadador. La distribución de sesiones es marcadamente desigual: un único nadador concentra el 37,2 % de los registros, tres nadadores acumulan el 67 %, y cuatro aportan una sola sesión cada uno (Figura 3.2). Esta concentración implica que un modelo entrenado sobre estos datos aprendería esencialmente el patrón de uno o dos atletas, comprometiendo la generalización entre individuos e introduciendo el riesgo de memorizar sujetos en lugar de aprender relaciones transferibles.

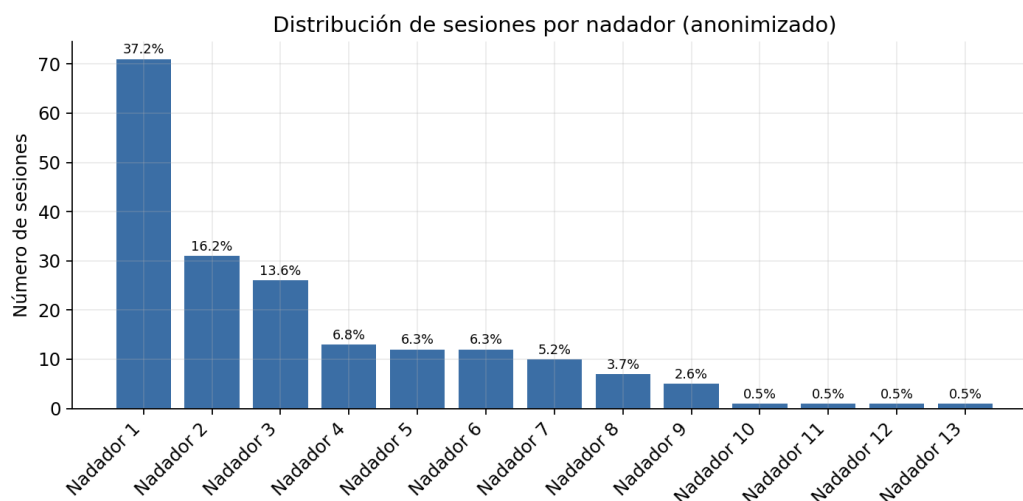


Figura 3.2: Distribución de sesiones por nadador (etiquetas anonimizadas). La concentración en unos pocos atletas limita la capacidad de generalización inter-individual.

Señal bienestar-rendimiento en datos reales. Sobre las sesiones con desenlace válido, la correlación entre el índice de bienestar y la diferencia de tiempo es débil ($r \approx -0,18$, $n = 96$). La dirección es la esperada, ya que mejor bienestar asociado a mejores tiempos, pero la magnitud es modesta y estadísticamente frágil dado el tamaño muestral. Este resultado es importante para acotar las expectativas: la relación que motiva el índice de bienestar existe en los datos reales de forma tenue, no como la asociación fuerte que más adelante se observará en el conjunto sintético.

Cobertura de categóricas. Las sesiones de baja intensidad predominan: R1 y R1-2 representan en conjunto cerca del 77% de los registros, sin sesiones de tipo TL. El estado de ánimo se concentra en *normal* (55,0%), con un 6,3% de reportes *muy mal*, y el dolor aparece en el 26,2% de las sesiones. La cobertura desigual de tipos de sesión e intensidades restringe la variedad de contextos de entrenamiento que un modelo podría aprender a distinguir.

En síntesis, el corpus real es reducido, está fuertemente desbalanceado por nadador y presenta el desenlace cronométrico ausente o corrupto en más de la mitad de los registros. Estas condiciones, propias de un proceso de captura incipiente, constituyen una base insuficiente para entrenar directamente un modelo supervisado generalizable, y son las que motivan la transición a la Fase 2.

3.1.5. De la Fase 1 a la Fase 2: motivación de la transición

Tras la visita a Fabio Toro Academy y la evaluación del prototipo inicial, la decisión de complementar los datos reales con un conjunto sintético se fundamentó en tres condiciones, que la caracterización anterior permite ahora sustentar cuantitativamente:

1. **Volumen y calidad insuficientes para el aprendizaje.** El desenlace cronométrico confiable se reduce a un centenar de sesiones y la información se concentra en pocos atletas, lo

que impide estimar relaciones que generalicen más allá de los sujetos sobrerrepresentados.

2. **Reporte diario poco sostenible.** La adherencia al formulario por parte de atletas juveniles resultó irregular, lo que se refleja en la proporción de campos ausentes y limita la posibilidad de construir series temporales completas por nadador.
3. **Alcance del prototipo inicial.** El modelo de la primera iteración clasificaba el rendimiento únicamente después de la sesión, sin capacidad de alertamiento anticipado, y dejaba sin cumplir uno de los objetivos específicos del trabajo.

En consecuencia y tras recibir asesoramiento de un experto en IA, se construyó un conjunto de datos sintético controlado que, *sin reemplazar los datos reales*, habilitara la prueba del pipeline completo, preprocesamiento, entrenamiento, validación y despliegue, en un entorno reproducible. Es importante señalar que esta decisión es coherente con la afirmación, sostenida en el cierre del trabajo, de que los datos reales recolectados resultaron insuficientes para un entrenamiento directo: el conjunto sintético no se presenta como evidencia empírica, sino como un banco de pruebas metodológico que permite verificar la viabilidad técnica del sistema mientras se consolida la captura de datos reales.

3.2. Fase 2: construcción y caracterización del conjunto sintético

3.2.1. Motivación metodológica

A partir de las limitaciones de la Fase 1, la Fase 2 persigue un objetivo acotado y explícito: disponer de un conjunto de datos suficientemente grande, estructurado y coherente con la lógica deportiva para ejercitar de extremo a extremo el pipeline analítico, equilibrando la fidelidad al dominio de los catorce nadadores de referencia. El valor de la Fase 2 es, por tanto, de validación técnica y no de descubrimiento empírico.

3.2.2. Supuestos y diseño de la generación

La generación se implementó como un pipeline de *generación basada en reglas* en Google Colab. Parte de los tiempos base de competencia de los catorce nadadores y de una cronología semanal anclada al macrociclo de diecisiete semanas de la academia (descrito en el Anexo C), con progresión de capacidades e incremento del volumen a lo largo de la temporada. Los supuestos mínimos, validados con el entrenador jefe, se declaran en el Anexo A; aquí interesa distinguir explícitamente qué relaciones fueron *inyectadas por diseño* y cuáles *emergen* del proceso, pues de ello depende la correcta interpretación del análisis posterior (Tabla 3.4).

Tabla 3.4: Relaciones inyectadas por diseño frente a propiedades emergentes en la generación sintética.

Inyectado por diseño	Emergente / estocástico
Asociación bienestar $\rightarrow$ tiempo real	Distribución final de la diferencia de tiempo
Efecto del dolor sobre estrés y motivación	Frecuencias observadas por estilo y tipo de sesión
Factores de intensidad por tipo de sesión	Casos atípicos ($\approx 5\%$ de contraste)
Tasa de asistencia ($\approx 92\%$) y ruido gaussiano	Valores únicos de tiempo real e índice de bienestar

Conviene anotar una precisión de honestidad metodológica: el dolor se genera de forma independiente entre sesiones, con una probabilidad cercana al 15% , y no como un proceso con persistencia del 40% entre días. Asimismo, el índice de bienestar de la generación es exactamente el definido en la Ecuación 3.3: la *misma* fórmula normalizada en $[0, 1]$ se aplica en ambas fases, por lo que los umbrales basados en él son directamente comparables entre la Fase 1 y la Fase 2.

3.2.3. Estructura general del conjunto generado

La Tabla 3.5 resume los indicadores del conjunto sintético.

Tabla 3.5: Indicadores estructurales del conjunto sintético (reproducibles con semilla 42).

Indicador	Valor
Registros totales	1.680
Columnas	17
Nadadores	14
Sesiones con asistencia efectiva	1.559
Tasa de asistencia	92,8 %
Diferencia media (real – planificado)	+1,41 s
Desviación estándar de la diferencia	2,85 s
Correlación bienestar vs. diferencia de tiempo	$r = -0,70$
Sesiones con dolor reportado	239 (15,3 %)

Los indicadores muestran variabilidad suficiente y ausencia de comportamiento determinista: la dispersión de la diferencia de tiempo y la diversidad de valores de tiempo real descartan una señal trivial, mientras que el tamaño del corpus y la tasa de asistencia ofrecen cobertura adecuada para ejercitar el modelado. La diferencia media positiva (+1,41 s) indica que, en promedio, el tiempo real generado es ligeramente superior al planificado; este valor no es una observación empírica, sino una consecuencia directa de los factores de intensidad por tipo de sesión y del ruido inyectado.

Debe insistirse en que se trata de un conjunto sintético, cuya señal fue construida bajo reglas controladas y desarrolladas en conjunto con el entrenador jefe de la academia y, por tanto, no

constituye evidencia empírica del rendimiento real de los nadadores.

3.2.4. Análisis exploratorio: verificación del generador

El análisis exploratorio del conjunto sintético debe leerse como una *verificación de que el generador reproduce la estructura diseñada*, no como el descubrimiento de patrones sobre nadadores reales. Bajo esa óptica, tres ejes confirman la coherencia del diseño:

- **Distribución por estilo:** predomina el estilo libre, acorde con la práctica de la academia, seguido por mariposa, pecho, espalda y combinado.
- **Distribución por tipo de sesión:** predominan R2 y VO2, consistente con la carga prevista por el macrociclo programado en las reglas de generación.
- **Variables de bienestar:** se concentran en estados normales, con escenarios contrastantes suficientes para discriminar contextos saludables y comprometidos.

El indicador más visible es la asociación inversa fuerte entre el índice de bienestar y la diferencia de tiempo ($r = -0,70$): a mejores condiciones físicas y emocionales, el tiempo real tiende a acercarse o a mejorar el planificado. Esta señal fue *inyectada* como parte del esquema de generación (Tabla 3.4); recuperarla en el análisis confirma que el generador funciona como se diseñó, pero no debe interpretarse como evidencia empírica de dicha relación.

3.2.5. Alcance y limitaciones de la Fase 2

El principal valor del conjunto sintético es permitir validar el funcionamiento técnico y metodológico del sistema antes de su aplicación con datos empíricos. En consecuencia, las conclusiones predictivas obtenidas sobre este conjunto no deben generalizarse como evidencia del rendimiento real de los nadadores.

El contraste entre ambas fases delimita con precisión este alcance: la asociación bienestar-rendimiento es de aproximadamente $r = -0,70$ en el conjunto sintético, donde fue inyectada, frente a $r \approx -0,18$ en los datos reales, estimada sobre un centenar de observaciones. Es decir, el conjunto sintético incorpora una relación sustancialmente más fuerte que la observada empíricamente; por ello, el desempeño que un modelo alcance sobre datos sintéticos sobreestima lo que cabría esperar con datos reales. La incorporación progresiva de datos reales en una fase posterior permitirá contrastar, debilitar o recalibrar los patrones aquí observados. Hasta entonces, los resultados deben leerse como una demostración de viabilidad metodológica y no como una caracterización empírica del grupo.

3.3. Trazabilidad de las variables y su rol en el modelado

Para preservar la coherencia entre la preparación de los datos y el modelado, conviene declarar explícitamente el rol de cada variable, distinguiendo las que constituyen *insumos predictores* de las

que *definen el desenlace* que el sistema busca anticipar (Tabla 3.6). Esta distinción es la salvaguarda principal frente a la fuga de información, fenómeno en el que una variable que forma parte de la definición del objetivo se utiliza, además, como predictora, induciendo una circularidad que infla artificialmente el desempeño aparente del modelo.

Tabla 3.6: Rol previsto de las variables en el modelado.

Rol	Variables
Predictoras (disponibles antes del desenlace)	Volumen, minutos por zona, tipo de sesión, RPE, horas de sueño, días desde competencia, fase del macrociclo, sRPE, ratio aeróbico.
Constitutivas del desenlace	Diferencia de tiempo (tiempo real vs. planificado) y los componentes del bienestar (ánimo, motivación, estrés, dolor) cuando se emplean para definir el bajo rendimiento.

3.4. Resumen del capítulo

Este capítulo desarrolló la Etapa 1 de la metodología: la identificación, recolección y preparación de los datos, organizada en dos fases iterativas y alineada con las fases de comprensión y preparación de datos del ciclo de vida estándar de un proyecto de aprendizaje automático. En la Fase 1 se establecieron los mecanismos de recolección mediante un formulario aplicado al finalizar cada sesión, se justificó la selección de variables desde una perspectiva multidisciplinaria y se documentó un pipeline de siete etapas con sus variables derivadas explícitas. La caracterización del corpus real, 191 sesiones de 13 nadadores, con el desenlace cronométrico ausente o corrupto en más de la mitad de los registros, fuerte concentración por atleta y una señal bienestar-rendimiento débil, mostró que los datos disponibles eran insuficientes para un entrenamiento directo, lo que motivó la Fase 2. En esta, se construyó y caracterizó un conjunto sintético controlado, cuyas propiedades se examinaron como verificación del generador y no como evidencia empírica. Finalmente, se declaró el rol de cada variable en el modelado.

Desarrollo del modelo analítico y predictivo para la detección temprana de bajo rendimiento deportivo

Esta sección describe el proceso de construcción, entrenamiento y evaluación del modelo analítico y predictivo. Su finalidad es detectar de manera anticipada patrones de bajo rendimiento en nadadores juveniles de Fabio Toro Academy a partir de los datos recolectados y caracterizados. La construcción se fundamenta en técnicas de aprendizaje automático supervisado, documentadas en la literatura como herramientas efectivas para la predicción de riesgos en el ámbito deportivo [Cordeiro et al., 2025](#); [Carvalho et al., 2024](#).

El desarrollo se realizó en dos fases iterativas. La **Fase 1** abordó la detección de bajo rendimiento como un problema de clasificación binaria, comparando tres algoritmos (Regresión Logística, Random Forest y XGBoost). La **Fase 2** amplió el alcance hacia una tarea de regresión: predecir el tiempo real de 100 m de la sesión siguiente mediante un *Multi-Layer Perceptron* (MLP). El detalle técnico ampliado de ambas fases se encuentra en el Anexo B.

4.1. Fase 1: Construcción y evaluación del modelo predictivo de bajo rendimiento

4.1.1. Definición del parámetro de riesgo: bajo rendimiento

La variable objetivo, denominada *bajo_rendimiento*, es una clasificación binaria: cada sesión registrada se etiqueta como **Normal (clase 0)** o **Bajo Rendimiento (clase 1)**. Su definición no se basa en un único indicador, sino en un enfoque multicriterio que refleja la naturaleza multidimensional del rendimiento deportivo juvenil, fenómeno que involucra factores fisiológicos, psicológicos y contextuales de manera simultánea [Bernaciková et al., 2022](#); [Qian and Lee, 2025](#). Se definieron siete criterios de riesgo (Cuadro 4.1).

Cada criterio se evalúa de forma binaria y se calcula un **puntaje de riesgo** (*risk_score*) como la suma de los criterios activos. Una sesión se clasifica como bajo rendimiento cuando el puntaje es igual o superior a 2, umbral establecido junto al cuerpo técnico bajo la premisa de que la coincidencia de dos o más señales simultáneas constituye un patrón significativo, mientras que un único factor puede ser circunstancial. La distribución resultante mostró un desbalance moderado

Tabla 4.1: Criterios de riesgo para la construcción de la variable objetivo *bajo_rendimiento*.

N.º	Criterio	Condición	Dimensión evaluada
1	Subrendimiento cronométrico	$\text{tiempo_real} > \text{tiempo_esperado}$	Rendimiento técnico
2	Esfuerzo percibido alto	$\text{RPE} \geq 7$	Carga percibida
3	Bienestar bajo	$\text{índice_bienestar} \leq 0,3$	Bienestar general
4	Motivación baja	$\text{motivación} \leq 4$	Estado psicológico
5	Estrés elevado	$\text{nivel_estrés} \geq 4$	Estado emocional
6	Presencia de dolor	$\text{dolor_presente} = 1$	Integridad física
7	Estado de ánimo muy malo	$\text{estado_ánimo} = 1$ (Muy mal)	Bienestar emocional

(aproximadamente 64 % Normal y 36 % Bajo Rendimiento), frecuente en problemas de detección de riesgo [Eetvelde et al., 2021](#) y atendido mediante técnicas de balanceo de clases.

El 64 % / 36 % reportado se obtuvo con los datos y con la versión del índice de bienestar disponibles cuando se entrenó este modelo; durante la recolección de datos se adoptó después una versión normalizada y unificada de ese índice, con la que el criterio «índice de bienestar $\leq 0,3$ » se activa menos veces y la proporción de sesiones de bajo rendimiento baja (en torno a 70/30). Por ello, volver a entrenar el modelo con el índice unificado y con un corpus más amplio forma parte de la recalibración prevista. Varias de las variables que definen el «bajo rendimiento» (por ejemplo, el dolor, el índice de bienestar o el RPE) también se entregan al modelo como variables de entrada; esto puede inflar sus resultados, y el problema y su tratamiento de este fenómeno se detalla más adelante.

4.1.2. Modelos de clasificación seleccionados

Se seleccionaron tres algoritmos de clasificación supervisada de familias complementarias: Regresión Logística (modelo lineal interpretable), Random Forest (ensamble de *bagging*) y XGBoost (ensamble de *boosting*). Esta elección sigue la recomendación metodológica de comparar algoritmos con distintos supuestos antes de seleccionar la mejor alternativa para un problema específico [Shetty, 2023](#); [Yang and Shami, 2020](#), y permite cubrir un espectro que va desde la interpretabilidad lineal hasta el alto rendimiento del *boosting*.

4.1.3. Proceso de entrenamiento y selección de hiperparámetros

El modelo de la Fase 1 se entrenó con datos reales: el subconjunto de sesiones con asistencia efectiva sobre las que es posible computar la etiqueta multicriterio de bajo rendimiento, que asciende a 140 sesiones. Se seleccionaron 18 variables predictoras que capturan las dimensiones de carga física, rendimiento cronométrico, percepción y bienestar, y contexto deportivo. Las variables categóricas se

4.1. Fase 1: Construcción y evaluación del modelo predictivo de bajo rendimiento 33

transformaron mediante codificación ordinal (*Label Encoding*) y los valores faltantes se imputaron con la mediana de cada variable. El conjunto se dividió mediante partición estratificada en 80 % para entrenamiento (112 sesiones) y 20 % para prueba (28 sesiones), preservando la proporción de clases; el conjunto de prueba resultante contiene 28 sesiones (18 Normal y 10 Bajo Rendimiento).

Los hiperparámetros se optimizaron mediante **GridSearchCV** con **validación cruzada estratificada de 5 folds**, utilizando el F1-Score como métrica de optimización en lugar de la exactitud, para balancear la detección de nadadores en riesgo (*recall*) con la precisión de las alertas y evitar el sesgo hacia la clase mayoritaria [Chicco and Jurman, 2020](#); [Bischl et al., 2023](#). La Regresión Logística se entrenó con estandarización previa (*StandardScaler*); el desbalance se compensó con `class_weight` y, en XGBoost, con `scale_pos_weight` y *early stopping*. Debe advertirse, no obstante, que este conjunto de 18 predictoras incluye variables que participan en la definición del objetivo (por ejemplo, `diferencia_tiempo`, `indice_bienestar`, `dolor_presente` o `rpe`), lo que introduce una fuga de información cuya influencia se discute mas adelante.

4.1.4. Resultados obtenidos y análisis

Todas las métricas reportadas corresponden al conjunto de prueba, no observado durante el entrenamiento, lo que garantiza una evaluación imparcial de la capacidad de generalización. El Cuadro 4.2 sintetiza el desempeño de los tres modelos.

Tabla 4.2: Métricas comparativas de los tres modelos en el conjunto de prueba.

Modelo	Accuracy	Precision	Recall	F1-Score	AUC-ROC	CV F1 (μ)
Regresión Logística	0.7857	0.7500	0.6000	0.6667	0.9167	0.9029
Random Forest	0.7500	0.7143	0.5000	0.5882	0.9167	0.9114
XGBoost	0.8214	0.7778	0.7000	0.7368	0.8694	0.9114

XGBoost obtuvo el mejor equilibrio entre detección y confiabilidad de las alertas: alcanzó el mayor F1-Score (0.7368) y el mayor *recall* (0.700), detectando 7 de los 10 nadadores en riesgo, frente a 6 de la Regresión Logística y 5 de Random Forest. La Regresión Logística y Random Forest presentan un AUC-ROC ligeramente superior (0.917 frente a 0.869), lo que refleja probabilidades mejor calibradas, aunque XGBoost genera mejores decisiones finales con el umbral estándar de 0.5.

Lectura operativa de las métricas. Más allá de las cifras, interesa traducirlas a consecuencias para la academia. Un *recall* de 0.70 significa que, en el conjunto evaluado, 3 de cada 10 sesiones de bajo rendimiento pasan inadvertidas (falsos negativos); una precisión de 0.78 implica que, de cada diez alertas emitidas, alrededor de ocho corresponden a sesiones efectivamente en riesgo y dos son falsas alarmas. En un sistema de alerta temprana orientado a proteger al atleta, el costo de un falso negativo (un nadador en riesgo que no recibe intervención) es asimétricamente mayor que el de un falso positivo (una revisión innecesaria por parte del cuerpo técnico). Esa asimetría es la que justifica haber optimizado el F1-Score y haber priorizado el *recall* en la selección del modelo:

la elección de XGBoost no responde solo a una cifra más alta, sino a una decisión explícita sobre qué tipo de error resulta más tolerable en este contexto.

Política de umbral derivada de la curva Precision-Recall. La curva Precision-Recall (disponible en el Anexo B) muestra que XGBoost mantiene una precisión perfecta hasta un *recall* cercano a 0.5; es decir, sus predicciones de mayor confianza no producen falsas alarmas. Esta lectura habilita una política operativa de dos niveles, derivable directamente de los datos ya disponibles: las alertas situadas en el rango de alta confianza, donde la precisión es máxima, pueden tratarse como *alertas de alta certeza* que justifican una intervención directa, mientras que las que aparecen al extender la detección hacia *recalls* mayores, donde la precisión decrece, se tratarían como *señales para revisar*, sujetas al juicio del entrenador. De este modo, la plataforma no necesita limitarse a una alerta binaria: puede graduar la respuesta según la confianza de la predicción, lo que constituye una implicación práctica concreta de la evaluación realizada.

Importancia de variables y fuga de información. En XGBoost, la presencia de dolor (ganancia 0.561) y el índice de bienestar (0.216) concentran el 77.7 % de la importancia, con un patrón coherente con el de Random Forest y con los signos de los coeficientes de la Regresión Logística, donde el dolor y el estrés operan como factores de riesgo y el bienestar y la motivación como factores protectores. Ahora bien, esta concentración debe interpretarse con cautela. Varias de las variables más influyentes (dolor, índice de bienestar, motivación, estrés, estado de ánimo) forman parte de la definición misma del objetivo (Cuadro 4.1), de modo que el modelo dispone, entre sus predictoras, de información que constituye la etiqueta. Esta fuga de información, infla el desempeño aparente y explica buena parte de la dominancia observada: el modelo no solo predice el bajo rendimiento, en parte lo reconstruye. En consecuencia, la importancia de variables aquí reportada debe leerse como una *hipótesis a verificar* y no como una conclusión.

4.1.5. Conclusión y selección del mejor modelo

El análisis permite concluir que **XGBoost es el modelo más adecuado** para su integración en la plataforma, por su mayor F1-Score (0.7368) y *recall* (0.700), criterio determinante en un sistema de alerta temprana orientado a la protección del atleta, sin sacrificar exactitud (0.8214) ni precisión (0.7778), con un sobreajuste controlado. Aunque la Regresión Logística ofrece mayor interpretabilidad y un AUC ligeramente superior, su menor *recall* (0.600) implicaría no detectar a 4 de cada 10 nadadores en riesgo, un compromiso poco aceptable para este propósito. Esta selección se sostiene como decisión de ingeniería para el prototipo.

Finalmente, el conjunto de datos disponible, aunque representativo del caso de estudio, es de tamaño moderado, lo que limita la precisión de las estimaciones y no abarca todos las casuísticas que se pueden presentar a lo largo de la temporada; a medida que la academia alimente la plataforma con nuevos registros, será posible reentrenar el modelo, eliminar las variables constitutivas del objetivo y reducir la brecha entre la validación cruzada y el conjunto de prueba, proceso que la arquitectura modular facilita.

4.2. Fase 2: Predicción del rendimiento futuro mediante un modelo MLP

Mientras la Fase 1 abordó la detección de bajo rendimiento como clasificación binaria, la Fase 2 amplía el alcance hacia una tarea de regresión: estimar el tiempo real esperado en 100 m para la sesión siguiente. Este giro ofrece al cuerpo técnico una salida más granular, expresada directamente en segundos, útil para anticipar desviaciones de rendimiento y ajustar la carga del día siguiente.

4.2.1. Motivación metodológica de la Fase 2

El objetivo concreto es estimar el tiempo real de 100 m en la sesión $t+1$ utilizando exclusivamente información disponible al cierre del día actual t . Esta restricción es deliberada: garantiza que el modelo opere en condiciones realistas de despliegue, sin recurrir a variables futuras ni a información que produzca *leakage* [Bischi et al., 2023](#). A diferencia de la Fase 1, aquí la fuga de información se trató de forma explícita desde el diseño, excluyendo las variables que reconstruyen el objetivo. El modelo no reemplaza el criterio del entrenador, sino que actúa como herramienta de apoyo para identificar patrones de evolución, estabilidad o deterioro, traduciendo su salida a un lenguaje familiar para el dominio deportivo.

4.2.2. Preprocesamiento y selección de variables

Partiendo de 1680 registros, se eliminaron 121 filas de inasistencia y se verificó la ausencia de duplicados exactos y de inconsistencias lógicas; tras construir las variables temporales necesarias para predecir $t + 1$, se conservaron 1519 filas útiles. La partición se realizó mediante un *split* temporal estricto por fechas globales (70 % entrenamiento, 15 % validación, 15 % prueba), y la estandarización (*StandardScaler*) se ajustó solo sobre el entrenamiento, evitando la fuga de estadísticas futuras [Bischi et al., 2023](#). Se excluyeron además las columnas `diferencia_tiempo`, `srpe` e `indice_bienestar` por fuga de información o redundancia.

Las variables predictoras combinan información numérica de rendimiento, percepción y carga disponible al cierre del día t ; variables temporales derivadas (*lags* y ventanas móviles) que capturan tendencias recientes; y variables categóricas codificadas con `LabelEncoder`. Se identificó una correlación de 0.97 entre `tiempo_100m_plan_sec` y el *target*, que debe interpretarse de manera doble: como fortaleza, porque el tiempo planificado es altamente predictivo; y como limitación, porque su dominancia puede reducir la visibilidad del aporte de las variables de bienestar y carga.

4.2.3. Selección y arquitectura del modelo

Se seleccionó un *Multi-Layer Perceptron* (MLP) denso, descartando arquitecturas recurrentes (LSTM, GRU) porque el volumen disponible, de 1519 filas distribuidas en 40 combinaciones de nadador y estilo, es insuficiente para entrenar secuencias sin un riesgo elevado de memorización; además, los *lags* y ventanas móviles ya codifican la dependencia temporal relevante, y el MLP

simplifica la inferencia en producción. La arquitectura final está compuesta por tres capas ocultas densas (64, 32 y 16 neuronas con activación ReLU), con regularización L2 y *dropout* para mitigar el sobreajuste, y una capa de salida lineal de una sola neurona, coherente con la regresión sobre el tiempo en segundos. Se empleó MAE como función de pérdida, interpretable directamente en segundos por los entrenadores [Chicco and Jurman, 2020](#), junto con el optimizador Adam y mecanismos de *early stopping*, reducción de la tasa de aprendizaje y conservación del mejor modelo.

4.2.4. Resultados del entrenamiento y evaluación

El Cuadro 4.3 resume el desempeño del modelo en los tres conjuntos.

Tabla 4.3: Desempeño del modelo MLP en entrenamiento, validación y prueba.

Métrica	Entrenamiento	Validación	Prueba
MAE (seg)	2.468	3.223	3.688
RMSE (seg)	3.108	4.155	4.547
R^2	0.9030	0.8562	0.7780

El MAE de prueba de 3.688 s representa un error operativo moderado para una predicción de corto plazo en 100 m en un contexto formativo, y el R^2 de 0.7780 indica que el modelo explica una proporción considerable de la variabilidad fuera de la muestra. La diferencia de aproximadamente 1.22 s en el MAE entre entrenamiento y prueba sugiere un sobreajuste leve a moderado, coherente con el tamaño limitado del *dataset* y la variabilidad entre nadadores. Dado que el conjunto es sintético, estas métricas describen la viabilidad técnica del enfoque y no su precisión definitiva en condiciones reales [Carvalho et al., 2024](#).

Sentido del error y del sesgo. La distribución de errores en prueba (Figura B.10 del Anexo B) presenta una media del error, calculado como valor predicho menos valor real, de aproximadamente +2,55 s; es decir, el modelo tiende a sobreestimar el tiempo y, por tanto, a pronosticar en promedio un rendimiento ligeramente peor del que ocurre. En términos operativos, este sesgo es *conservador* para un sistema de alerta, porque inclina las predicciones hacia el lado de la precaución y reduce el riesgo de pasar por alto un deterioro; sin embargo, un sesgo sistemático de esta magnitud podría inducir fatiga de alarmas si la plataforma reacciona ante cada sobreestimación. Por ello, la corrección de este sesgo es uno de los objetivos de la recalibración con datos reales, que permitiría centrar el error alrededor de cero sin sacrificar la sensibilidad del sistema.

4.2.5. Alcance y limitaciones de la Fase 2

El modelo MLP ofrece una primera aproximación útil para predecir el rendimiento futuro en segundos, transformando el sistema desde una alerta binaria hacia una salida continua interpretable directamente por el entrenador. No obstante, los resultados no deben tomarse como evidencia

definitiva del rendimiento real de los nadadores, porque el *dataset* sigue siendo sintético y parte de la señal fue diseñada de forma controlada. La alta correlación entre el tiempo planificado y el *target* (0.97) puede favorecer el rendimiento predictivo agregado, pero también ocultar el aporte específico de las variables de bienestar y carga, verificable solo con datos empíricos. El tamaño del conjunto y la disponibilidad desigual de registros por combinación de nadador y estilo limitan la robustez del modelo y son consistentes con el sobreajuste observado; la incorporación progresiva de datos reales permitirá recalibrarlo, corregir el sesgo de sobreestimación y mejorar su estabilidad. Aun así, la Fase 2 aporta un avance metodológico relevante: una predicción continua, más informativa para la planificación diaria y compatible con la arquitectura modular de la plataforma.

4.3. Resumen del capítulo

Este capítulo describió la construcción de los modelos analíticos y predictivos que sustentan la plataforma propuesta, organizada en dos fases iterativas. En la **Fase 1** se abordó la detección de bajo rendimiento como clasificación binaria, a partir de una variable objetivo multicriterio construida con el cuerpo técnico, comparando tres algoritmos (Regresión Logística, Random Forest y XGBoost) mediante validación cruzada estratificada y búsqueda de hiperparámetros; el análisis seleccionó a XGBoost por su mejor equilibrio entre precisión y sensibilidad. Esa selección se acompaña de una lectura operativa de las métricas (costo de falsos negativos frente a falsos positivos y una política de alertas de dos niveles) y del reconocimiento explícito de una fuga de información: como varias de las variables más influyentes constituyen la etiqueta, la dominancia del dolor y el bienestar se presenta como una hipótesis a verificar con datos reales y sin esas variables, y no como una conclusión. En la **Fase 2** se amplió el alcance hacia una predicción continua del tiempo real de 100 m de la sesión siguiente, mediante un MLP entrenado sobre el *dataset* sintético depurado, con tratamiento explícito de la fuga, *features* temporales y una arquitectura ajustada al tamaño del conjunto; sus resultados, expresados en segundos, ofrecen una salida más granular y operativamente útil, con un sesgo conservador de sobreestimación que la recalibración futura deberá corregir. En conjunto, ambas fases configuran un sistema predictivo coherente, trazable y modular, base para la integración del modelo en la plataforma web y su posterior recalibración con datos empíricos.

Diseño y desarrollo del prototipo de la plataforma web modular

Esta sección presenta las decisiones de diseño y desarrollo del prototipo de la plataforma web modular propuesta, cuyo propósito consistió en integrar la centralización de datos, el análisis y la generación de alertas tempranas para el monitoreo del rendimiento deportivo en nadadores juveniles. La construcción del prototipo se fundamentó en principios de arquitectura modular, escalabilidad y separación de responsabilidades, con el fin de garantizar la sostenibilidad técnica del sistema y su capacidad de evolución futura.

La plataforma fue concebida como un sistema web orientado a múltiples actores, entre ellos entrenadores, deportistas y acudientes, permitiendo la gestión, visualización y comunicación eficiente de información relevante del proceso formativo. Su diseño respondió a las necesidades identificadas en el caso de estudio Fabio Toro Academy, en el que la información de rendimiento se encontraba previamente dispersa en múltiples formatos y sin mecanismos automatizados de análisis ni de alerta temprana.

La sección se estructuró en apartados que abordaron, de manera progresiva, los requerimientos del sistema, la arquitectura propuesta, el diseño de los módulos principales, la integración del modelo analítico y el desarrollo del prototipo funcional, así como las fases previstas para su crecimiento.

5.1. Requerimientos del sistema

Con base en el análisis del problema, los objetivos del proyecto y las características del caso de estudio, se definieron los requerimientos funcionales y no funcionales que orientaron el diseño y la implementación del prototipo de la plataforma web modular.

Requerimientos funcionales Los requerimientos funcionales describieron las capacidades que debía ofrecer la plataforma para apoyar el monitoreo del rendimiento deportivo juvenil en el contexto del caso de estudio.

Centralización de datos de rendimiento

- **RF_01. Registro de información deportiva:** La plataforma debe permitir el registro estructurado de información asociada al proceso formativo de los nadadores juveniles, inclu-

yendo cargas de entrenamiento, indicadores fisiológicos, variables de bienestar subjetivo y registros de asistencia.

- **RF_02. Almacenamiento centralizado de datos:** El sistema debe almacenar de manera centralizada la información registrada, garantizando su disponibilidad para consulta, análisis y visualización desde una única fuente de datos.
- **RF_03. Consulta histórica de registros:** La plataforma debe permitir la consulta del histórico de datos de cada nadador, facilitando el seguimiento de su evolución deportiva.
- **RF_04. Actualización de registros:** El sistema debe permitir la actualización de la información previamente registrada, con el propósito de corregir inconsistencias o incorporar nuevos datos relevantes al seguimiento deportivo.
- **RF_05. Organización de datos por nadador y período:** La plataforma debe permitir la organización y recuperación de la información por nadador, fecha, período de entrenamiento o tipo de medición, favoreciendo su análisis posterior.

Gestión de perfiles de usuario

- **RF_06. Administración de perfiles de acceso:** El sistema debe contemplar perfiles de usuario diferenciados para entrenadores, deportistas y acudientes, de acuerdo con las necesidades y responsabilidades de cada actor dentro del proceso formativo.
- **RF_07. Autenticación de usuarios:** La plataforma debe permitir el acceso autenticado de los usuarios, asegurando que únicamente personas autorizadas puedan ingresar al sistema.
- **RF_08. Autorización basada en roles:** El sistema debe restringir el acceso a funcionalidades y datos según el perfil asignado, garantizando que cada usuario visualice únicamente la información correspondiente a su rol.
- **RF_09. Visualización personalizada por rol:** La plataforma debe presentar vistas, módulos e información adaptados al tipo de usuario, de modo que entrenadores, deportistas y acudientes accedan a contenidos pertinentes para su participación en el proceso formativo.

Integración del modelo analítico y predictivo

- **RF_10. Procesamiento de datos:** La plataforma debe procesar la información registrada para alimentar el modelo analítico y predictivo integrado en el sistema.
- **RF_11. Ejecución del modelo predictivo:** El sistema debe ejecutar el modelo analítico y predictivo sobre los datos disponibles, con el fin de identificar patrones asociados al bajo rendimiento deportivo.

- **RF_12. Generación de resultados predictivos:** La plataforma debe generar resultados relacionados con el nivel de riesgo de bajo rendimiento, según las variables disponibles.
- **RF_13. Disponibilidad de resultados para consulta:** El sistema debe almacenar y poner a disposición de los usuarios autorizados los resultados generados por el modelo, permitiendo su consulta posterior dentro de la plataforma.

Visualización de indicadores y resultados

- **RF_14. Presentación de indicadores clave de rendimiento:** La plataforma debe mostrar indicadores clave relacionados con el rendimiento deportivo de los nadadores, a partir de la información consolidada y procesada por el sistema.
- **RF_15. Visualización mediante dashboards interactivos:** El sistema debe ofrecer tableros interactivos para la representación gráfica de la información, facilitando la interpretación de tendencias, variaciones y comportamientos relevantes.
- **RF_16. Consulta segmentada de información:** La plataforma debe permitir la consulta segmentada de indicadores por nadador, período, tipo de sesión, con el propósito de facilitar la exploración de la información.
- **RF_17. Presentación comprensible de resultados:** El sistema debe presentar la información de manera clara y comprensible, considerando que los usuarios finales no necesariamente cuentan con formación técnica especializada.
- **RF_18. Detección de condiciones de riesgo:** La plataforma debe identificar condiciones asociadas al bajo rendimiento deportivo a partir de las salidas del modelo analítico y predictivo.
- **RF_19. Emisión automática de alertas:** El sistema debe generar alertas automáticas cuando se detectan patrones o comportamientos que indicaran posibles riesgos en el rendimiento del nadador.
- **RF_20. Asociación de alertas a nadadores específicos:** Cada alerta generada debe vincularse al deportista correspondiente, incluyendo el contexto necesario para su adecuada interpretación por parte del entrenador.
- **RF_21. Consulta de alertas históricas:** La plataforma debe permitir la consulta de alertas emitidas previamente, con el fin de apoyar el seguimiento de casos y la toma de decisiones informadas.

Soporte a la comunicación entre actores

- **RF_22. Acceso compartido a información relevante:** La plataforma debe facilitar el acceso compartido a información del proceso formativo entre entrenadores, deportistas y acudientes, de acuerdo con los permisos definidos para cada rol.
- **RF_23. Disponibilidad de resultados de seguimiento:** El sistema debe poner a disposición de los actores autorizados los indicadores, reportes y resultados relevantes del monitoreo deportivo, fortaleciendo la transparencia del proceso.
- **RF_24. Apoyo a la toma de decisiones:** La plataforma debe proporcionar información consolidada y contextualizada que brinde apoyo para la toma de decisiones relacionadas con entrenamiento, seguimiento y acompañamiento del nadador.

Requerimientos no funcionales Los requerimientos no funcionales establecen criterios de calidad que garantizan la correcta operación y sostenibilidad del sistema:

- **Modularidad:** La arquitectura debe permitir la incorporación, modificación o reemplazo de módulos sin afectar la estabilidad global del sistema.
- **Escalabilidad:** El prototipo debe estar preparado para un incremento futuro en el volumen de datos, número de usuarios o funcionalidades, sin degradar significativamente su desempeño.
- **Usabilidad:** La interfaz debe ser intuitiva y accesible, considerando que los usuarios finales no necesariamente poseen formación técnica especializada.
- **Seguridad y privacidad:** El sistema debe proteger la información sensible de los deportistas, cumpliendo con la normativa colombiana vigente en materia de protección de datos personales.
- **Mantenibilidad:** El diseño del software debe facilitar la corrección de errores, la actualización de componentes y la evolución del sistema a largo plazo.

5.1.1. Diagrama de casos de uso

Con el propósito de representar de manera gráfica las principales interacciones entre los actores y la plataforma propuesta, se elaboró un diagrama de casos de uso basado en los requerimientos funcionales identificados. Este diagrama permite visualizar las funcionalidades ofrecidas por el sistema desde la perspectiva de entrenadores, deportistas y acudientes, evidenciando cómo cada actor participa en los procesos de registro, consulta, seguimiento y monitoreo del rendimiento deportivo.

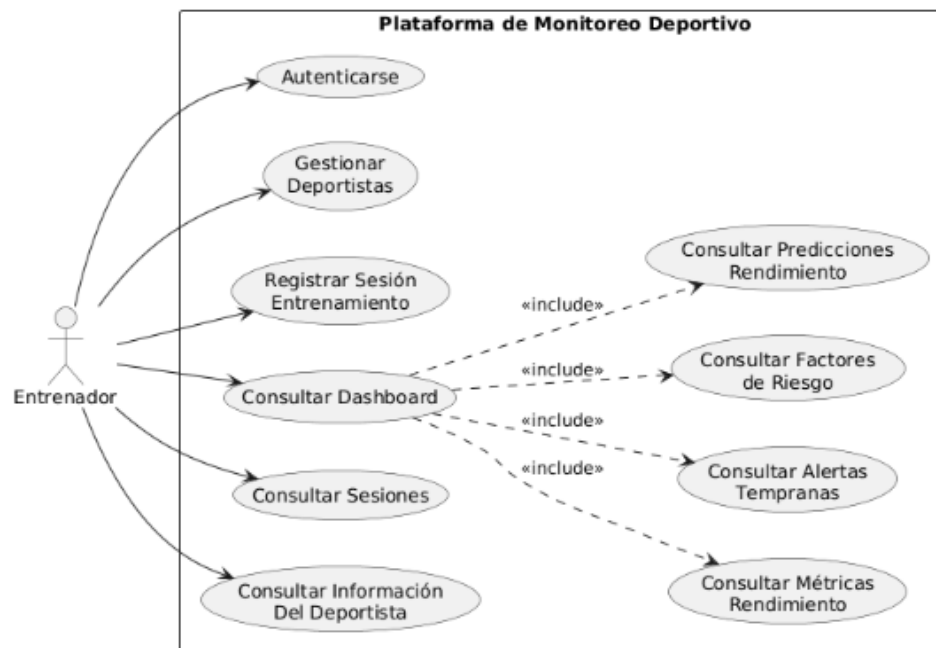


Figura 5.1: Diagrama caso de uso entrenador

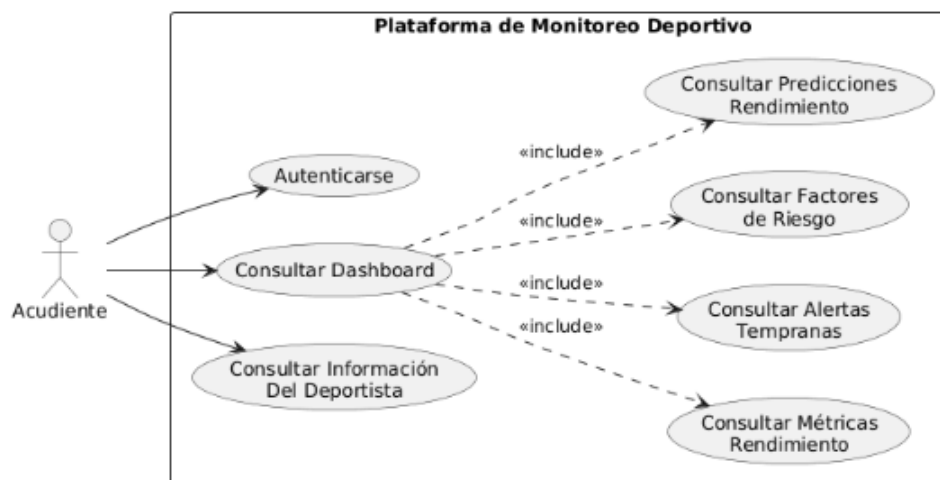


Figura 5.2: Diagrama caso de uso acudiente

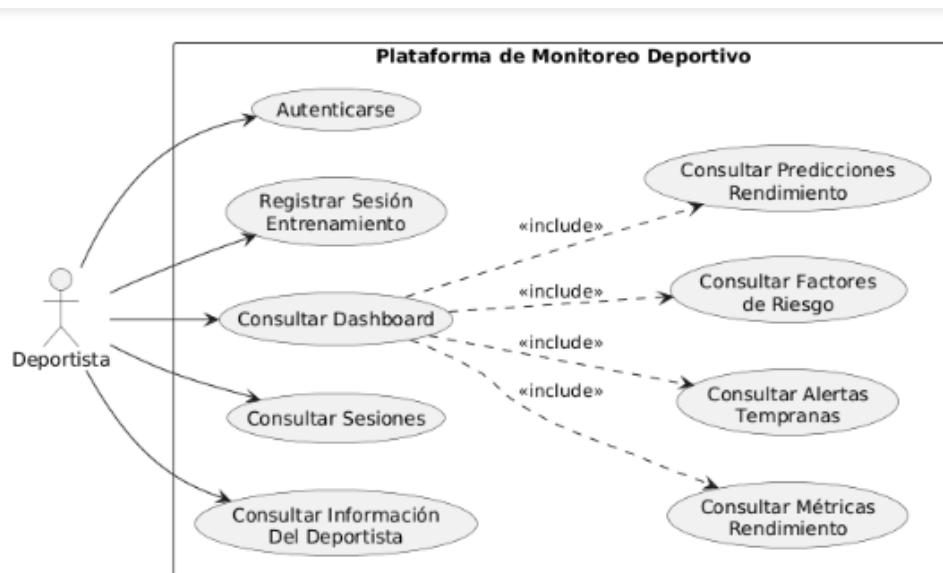


Figura 5.3: Diagrama caso de uso deportista

5.2. Arquitectura de la plataforma

La arquitectura de la plataforma web fue diseñada con un enfoque evolutivo, considerando tanto las necesidades inmediatas del caso de estudio Fabio Toro Academy como la proyección futura del sistema hacia una arquitectura modular. El diseño arquitectónico se plantea en fases progresivas, permitiendo validar funcionalidades con usuarios reales, recolectar retroalimentación y evolucionar hacia un sistema más escalable y desacoplado. A continuación, se describe el proceso del desarrollo arquitectónico.

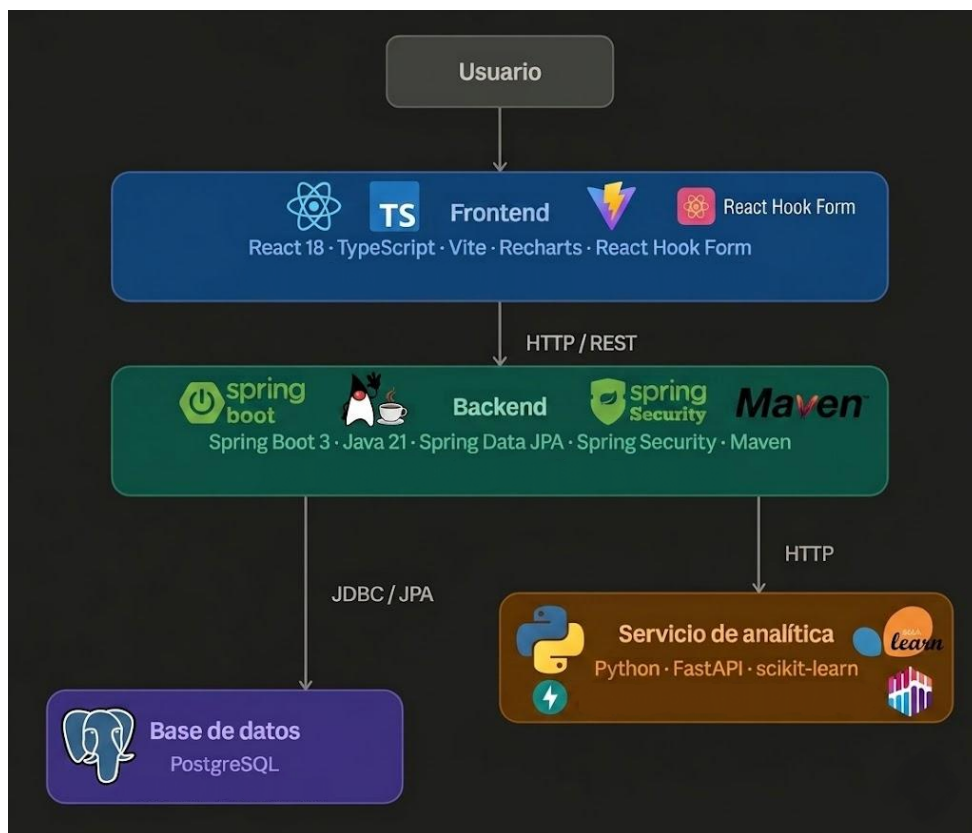


Figura 5.4: Diagrama arquitectura con stack tecnologico

5.2.1. Selección y justificación tecnológica

La selección del stack tecnológico empleado en el desarrollo del prototipo respondió a una combinación de criterios técnicos, académicos y contextuales. Por un lado, se consideró la naturaleza del problema abordado y las características funcionales y no funcionales definidas para el sistema; por otro, se tuvo en cuenta la experiencia previa del equipo de desarrollo en las tecnologías evaluadas, factor que resulta determinante en proyectos de investigación aplicada donde los tiempos de desarrollo son acotados y la curva de aprendizaje representa un costo significativo.

La selección no obedeció a criterios arbitrarios ni exclusivamente a preferencias personales, sino que se fundamentó en la alineación entre las capacidades ofrecidas por cada tecnología y los requerimientos específicos del sistema: modularidad, separación de responsabilidades, integración de componentes heterogéneos, escalabilidad progresiva y mantenibilidad a largo plazo. A continuación se presenta la justificación de cada decisión tecnológica relevante.

5.2.2. Frontend: React con TypeScript y Vite

Para el componente de presentación se seleccionó React 18 como biblioteca principal de construcción de interfaces, en conjunto con TypeScript como lenguaje de programación y Vite como herramienta de empaquetado y desarrollo. La elección de React se sustenta en su amplia adopción en el desarrollo de aplicaciones web de tipo SPA, su ecosistema maduro de componentes reutilizables y su modelo de programación declarativa basado en componentes, que favorece la separación de responsabilidades en la capa de presentación. TypeScript fue incorporado con el propósito de añadir tipado estático al desarrollo, reduciendo la probabilidad de errores en tiempo de ejecución y mejorando la mantenibilidad del código en el largo plazo, aspectos especialmente relevantes en sistemas que contemplan evolución por fases. Vite, por su parte, ofrece tiempos de compilación significativamente menores respecto a herramientas alternativas como Create React App, lo que optimiza el ciclo de desarrollo iterativo propio de un prototipo funcional.

La familiaridad del equipo con el paradigma de componentes de React, combinada con la disponibilidad de librerías especializadas para visualización de datos como Recharts y gestión de formularios como React Hook Form, consolidó esta elección como la más adecuada para las necesidades del proyecto. La arquitectura SPA resultante es coherente con los requerimientos de navegación fluida entre tableros de control, formularios de registro y vistas de detalle, sin recargas completas de página que pudieran afectar la experiencia del usuario.

5.2.3. Backend: Spring Boot con Java

Para el componente de lógica de negocio y persistencia se seleccionó Spring Boot 3 sobre Java 21. Esta decisión se sustentó en la robustez del framework para el desarrollo de APIs REST, su integración nativa con Spring Data JPA para el manejo de persistencia relacional, y su ecosistema de seguridad a través de Spring Security. Spring Boot ofrece un modelo de desarrollo orientado a convención sobre configuración, que reduce la complejidad inicial sin sacrificar la capacidad de personalización requerida en etapas más avanzadas del proyecto.

Java, como lenguaje de programación orientado a objetos ampliamente utilizado en entornos empresariales y académicos, brinda garantías de tipado estático, gestión de memoria y un ecosistema de herramientas de construcción en particular Maven que facilita la gestión de dependencias y la reproducibilidad del entorno de desarrollo. Así como la experiencia previa del equipo en el ecosistema Spring fue un factor determinante.

5.2.4. Servicio de analítica: Python con FastAPI y scikit-learn

Para el componente de analítica y predicción se seleccionó Python como lenguaje de implementación, FastAPI como framework de exposición del servicio y scikit-learn como biblioteca de aprendizaje automático. La elección de Python para este componente responde a su posición dominante en el ecosistema de ciencia de datos e inteligencia artificial, así como a la disponibilidad de bibliotecas especializadas Pandas, NumPy, scikit-learn, Keras, Joblib, Matplotlib y Seaborn que

cubren de manera integral las etapas del pipeline analítico: generación y procesamiento de datos, entrenamiento del modelo y evaluación de resultados.

5.2.5. Base de datos: PostgreSQL

Para la capa de persistencia se seleccionó PostgreSQL como sistema de gestión de bases de datos relacional. Esta elección se fundamenta en su madurez tecnológica, su soporte completo del estándar SQL, su robustez para el manejo de integridad referencial y su integración nativa con Spring Data JPA a través del dialecto Hibernate correspondiente. PostgreSQL ofrece, además, capacidades avanzadas de indexación y consulta que resultan relevantes para el crecimiento futuro del volumen de datos en el sistema. Su condición de software de código abierto, ampliamente adoptado tanto en entornos de investigación como en producción, lo posiciona como una opción técnicamente sólida y sostenible para el ciclo de vida del proyecto.

La naturaleza relacional del modelo de datos del sistema con entidades fuertemente relacionadas entre sí, como deportistas, sesiones de entrenamiento, métricas y predicciones justifica la preferencia por una base de datos relacional sobre alternativas NoSQL, cuyas ventajas son más pronunciadas en escenarios de datos no estructurados o de alta variabilidad esquemática.

5.2.6. Síntesis de la decisión tecnológica

En conjunto, el stack seleccionado React + TypeScript en el frontend, Spring Boot + Java en el backend, Python + FastAPI en el servicio de analítica y PostgreSQL como capa de persistencia constituye una combinación tecnológicamente coherente, orientada a la separación de responsabilidades entre componentes heterogéneos y respaldada tanto por criterios técnicos objetivos como por la experiencia acumulada del equipo de desarrollo. Esta combinación permitió avanzar con agilidad en la construcción del prototipo, sin comprometer la capacidad de evolución arquitectónica proyectada para fases posteriores.

5.3. Fase 1

5.3.1. Enfoque arquitectónico.

En esta primera fase del desarrollo, la plataforma adoptó una arquitectura centralizada fundamentada en un backend monolítico desarrollado, que expone servicios REST consumidos por una aplicación web de tipo SPA (Single Page Application). Este enfoque fue seleccionado por su simplicidad estructural, su facilidad de mantenimiento durante las etapas iniciales del proyecto y la reducción de la complejidad operativa requerida para las primeras interacciones con usuarios finales en el contexto del caso de estudio. La solución se complementa con un servicio externo de analítica y predicción, cuya responsabilidad comprende el procesamiento de datos ingresados en el registro de la sesión de entrenamiento, para brindar como respuesta la predicción de bajo rendimiento y los factores de riesgo asociados. De esta manera, si bien el núcleo del sistema se despliega como

un monolito, la integración con el componente analítico introduce desde esta fase una separación clara de responsabilidades entre la lógica de negocio central y el procesamiento avanzado de datos, sentando las bases para una evolución arquitectónica progresiva.

5.3.2. Arquitectura en capas del backend.

El backend de la plataforma sigue una arquitectura en capas (Layered / N-Tier), ampliamente utilizada en aplicaciones desarrolladas con Spring Boot y reconocida por favorecer la separación de responsabilidades, facilitar la mantenibilidad del sistema y establecer límites claros entre los distintos niveles de la solución. A partir de la revisión técnica del prototipo, se identifican las siguientes capas:

La capa de presentación está implementada mediante controladores REST responsables de exponer los endpoints HTTP, gestionar las solicitudes entrantes y documentar el contrato de la API a través de OpenAPI/Swagger. Esta capa incluye controladores orientados a la gestión de usuarios, consultas de listas de referencia, operaciones sobre tableros de control y el registro y detalle de sesiones de entrenamiento.

La capa de servicios o lógica de negocio concentra la lógica principal de la aplicación, incluyendo la gestión de usuarios, el registro de sesiones de entrenamiento, la orquestación de consultas analíticas mediante la invocación al servicio de ML, y la validación de reglas de negocio. Es en esta capa donde se materializa la integración con el componente de analítica: al registrar una sesión con asistencia confirmada, el servicio correspondiente construye la solicitud al modelo predictivo, invoca el servicio externo, procesa la respuesta y persiste la predicción, los factores de riesgo y las alertas generadas.

La capa de acceso a datos se implementa mediante repositorios JPA encargados de la persistencia y recuperación de información desde PostgreSQL, garantizando que la lógica de negocio permanezca desacoplada de los detalles de persistencia.

La capa de dominio y transferencia incluye las entidades JPA que representan el modelo de datos del sistema y los objetos de transferencia de datos utilizados tanto para la comunicación entre capas internas como para los contratos de entrada y salida de la API.

Este enfoque garantiza que los controladores no contengan lógica de negocio y que el acceso a datos esté correctamente encapsulado, lo cual resulta fundamental para la mantenibilidad y la evolución futura del sistema.

5.3.3. Modelo de Representación Arquitectónica C4

Para describir la arquitectura de la plataforma de forma estructurada y comprensible, se adopta el modelo C4, que permite representar el sistema desde distintos niveles de abstracción, desde el contexto más amplio hasta los componentes internos más relevantes.

En el Nivel de Contexto, independientemente de la fase, el contexto del sistema se mantiene constante, la plataforma se presenta como un sistema central de análisis de rendimiento deportivo

que interactúa con distintos actores entrenadores, deportistas, acudientes y administradores así como con sistemas externos, entre ellos el servicio de analítica predictiva encargado del procesamiento del modelo de aprendizaje automático.

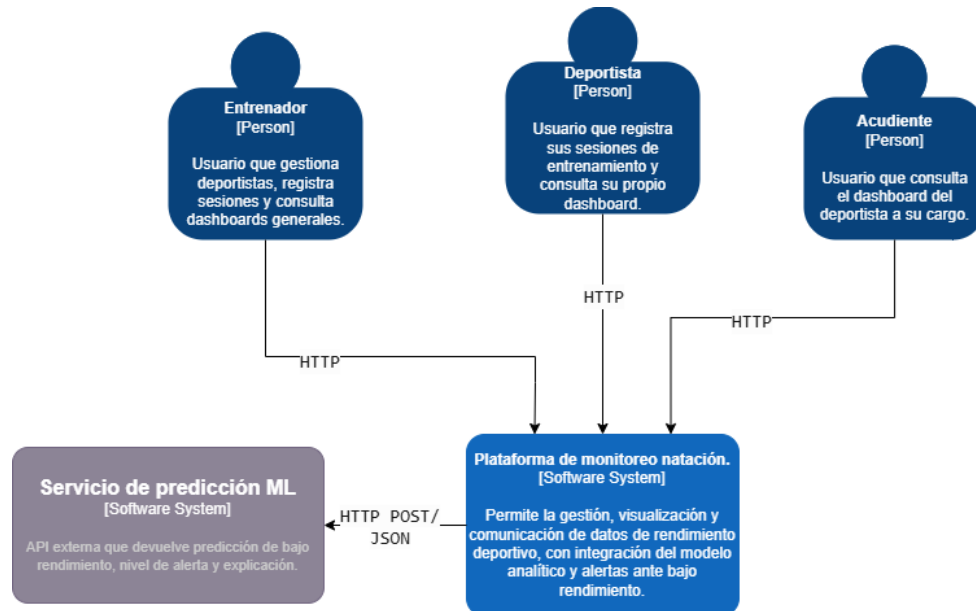


Figura 5.5: Diagrama contexto

En el Nivel de Contenedores, al igual que el de contexto independiente la fase a desarrollar se va a comportar de la misma manera, el sistema se compone de tres elementos principales: una aplicación web SPA que actúa como interfaz de usuario, un backend que centraliza la lógica de negocio y la persistencia, y un servicio externo de analítica que expone el modelo predictivo y por último la base de datos relacional. Cada contenedor cumple una responsabilidad claramente delimitada, y la comunicación entre ellos se realiza exclusivamente mediante interfaces HTTP con contratos definidos.

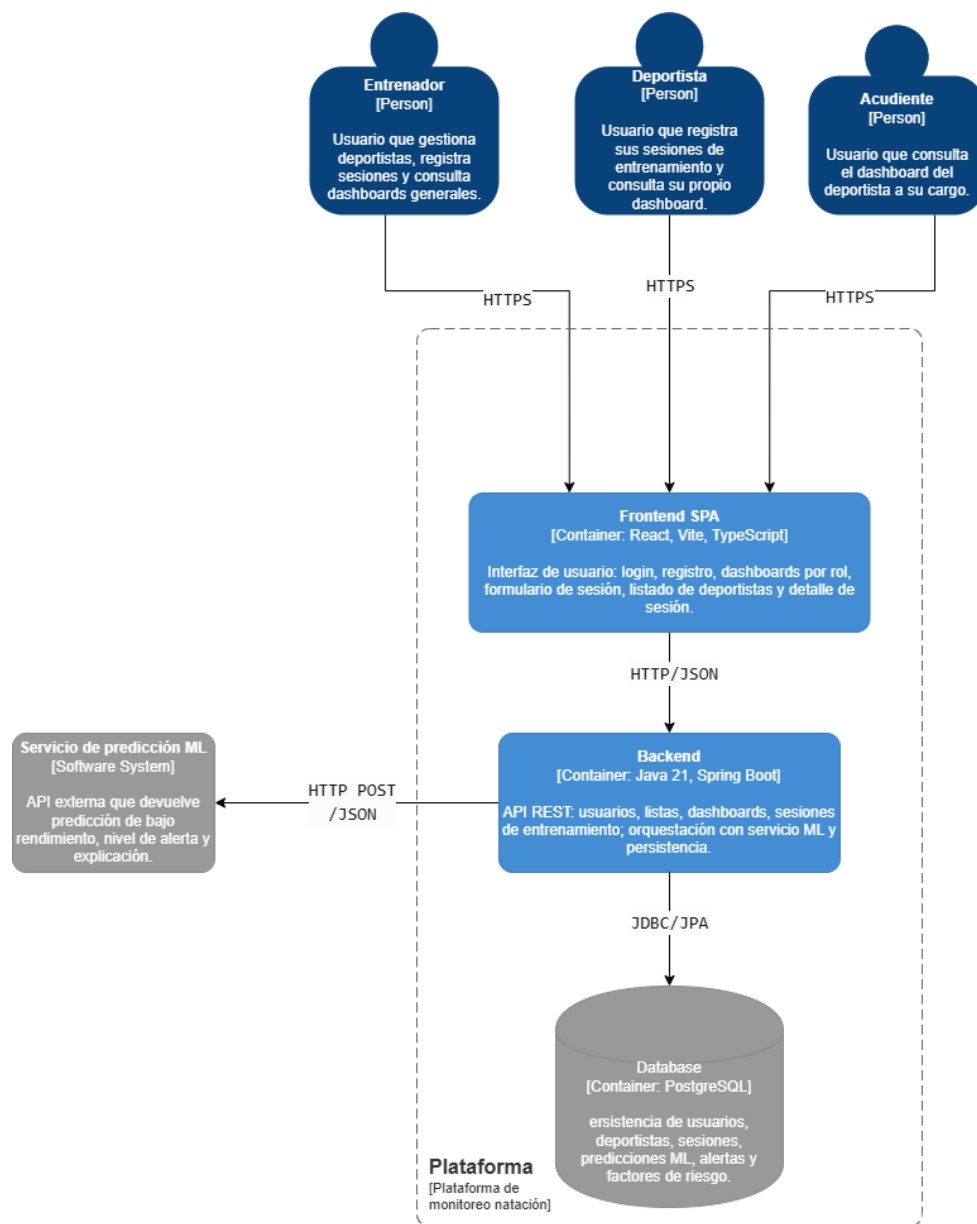


Figura 5.6: Diagrama contenedores

A nivel de componentes, el backend se encuentra organizado siguiendo una arquitectura en capas, donde cada componente asume responsabilidades específicas dentro del flujo de procesamiento de la aplicación. El diagrama presentado a continuación muestra los controladores encargados de la exposición de la API REST, los servicios que concentran la lógica de negocio, los componentes de integración utilizados para la comunicación con el servicio externo de analítica y los repositorios responsables de la persistencia de datos. Esta representación permite visualizar las dependencias

internas del sistema y comprender el flujo de información desde la recepción de una solicitud hasta el almacenamiento y consulta de los datos asociados.

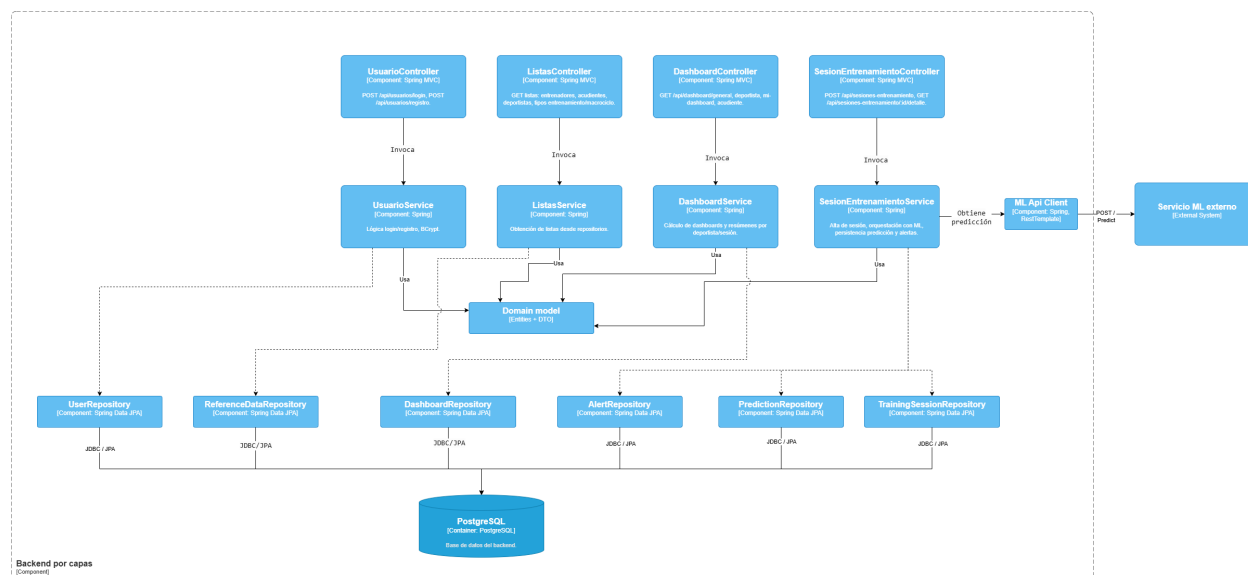


Figura 5.7: Diagrama componente backend - Fase 1

5.3.4. Definición Arquitectónica del Frontend

El componente de presentación corresponde a una SPA construida con React y Vite, organizada en páginas, servicios, componentes reutilizables y un estado global mínimo gestionado mediante React Context para la autenticación. La capa de servicios centraliza todas las llamadas al API REST del backend, manteniendo un único punto de configuración de la URL del servidor y un contrato claro con los endpoints disponibles. La gestión de formularios se apoya en React Hook Form, y la visualización de métricas y tendencias se implementa mediante Recharts. El enrutamiento entre vistas se realiza con React Router, con protección de rutas basada en el rol del usuario autenticado y un componente despachador que resuelve el tablero de control apropiado para cada perfil.

La autenticación del usuario se gestiona a través de un contexto global que persiste la sesión de forma local en el navegador, permitiendo la navegación protegida entre vistas sin necesidad de reautenticación en cada interacción. Todas las peticiones HTTP hacia el backend se canalizan a través de una función centralizada de fetch, lo que permite gestionar de manera uniforme la URL base del servidor y el tratamiento de errores de comunicación.

Las responsabilidades funcionales del frontend comprenden: autenticación y registro de usuarios con redirección por rol; registro de sesiones de entrenamiento mediante un formulario multietapa; visualización de tableros de control con filtros, tablas y gráficos de métricas de asistencia, volumen y probabilidad de bajo rendimiento; consulta del detalle de sesiones con visualización de alertas y

factores de riesgo; y listado de deportistas para el perfil de entrenador. Dando como resultado la siguiente interpretación a través del modelo C4.

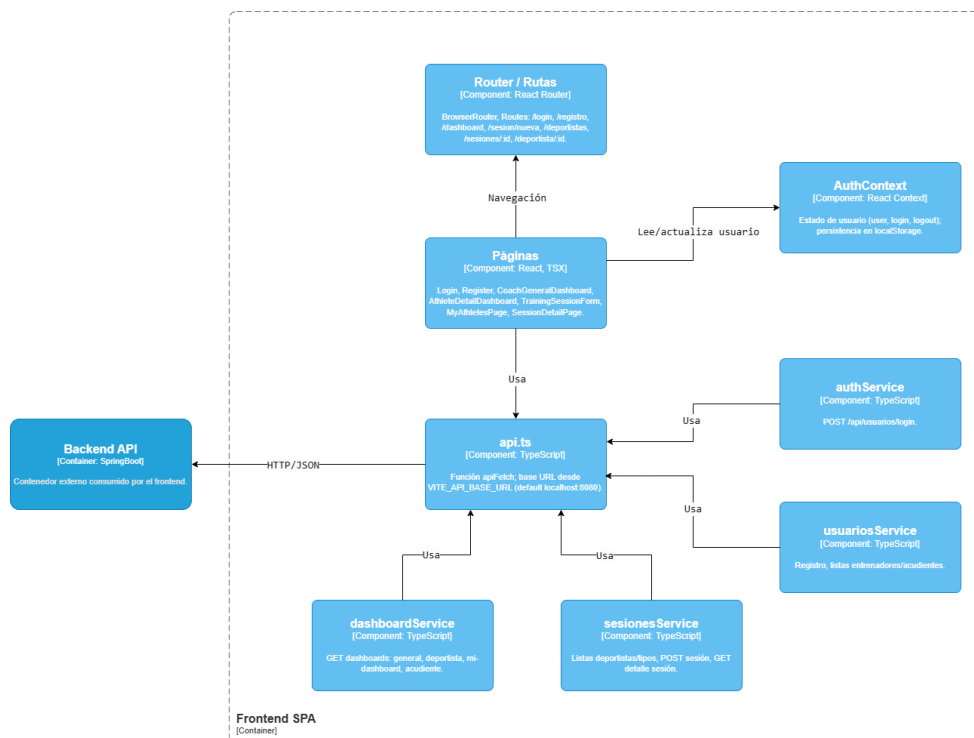


Figura 5.8: Diagrama componente frontend - Fase 1

5.3.5. Integración con el servicio de analítica y predicción

La plataforma integra un servicio externo de analítica y predicción desarrollado en Python con FastAPI, consumido por el backend mediante una invocación HTTP en el flujo de registro de sesión de entrenamiento. Cuando se registra una sesión con asistencia confirmada, el backend envía al endpoint `POST /predict` un objeto con las variables de carga de la sesión (volumen total y minutos en zona aeróbica y anaeróbica), los tiempos cronométricos planificado y real, el estado de bienestar y fisiológico del deportista (RPE, horas de sueño, estado de ánimo, nivel de estrés, motivación y presencia de dolor) y el contexto deportivo (días desde la última competencia y fase del macrociclo). Internamente, el servicio calcula las características derivadas requeridas por el modelo (sRPE, ratio aeróbico, índice de bienestar y diferencia entre tiempo real y planificado), evalúa los siete criterios de riesgo definidos en la construcción de la variable objetivo y ejecuta el mejor clasificador seleccionado durante el entrenamiento entre Regresión Logística, Random Forest y XGBoost según el F1-Score. La respuesta devuelve la predicción binaria de bajo rendimiento, la probabilidad asociada, el nivel categórico de riesgo en cuatro escalones (bajo, medio, alto y muy alto) determinado por umbrales sobre la probabilidad, y el detalle explicable de cada uno de los

siete factores evaluados con su valor actual y umbral de activación. Con base en esta información, el backend persiste la predicción y genera las entidades de alerta y factores de riesgo asociadas, que posteriormente se exponen al usuario a través de los tableros de control y las vistas de detalle de sesión. Dando como resultado la siguiente interpretación a través del modelo C4.

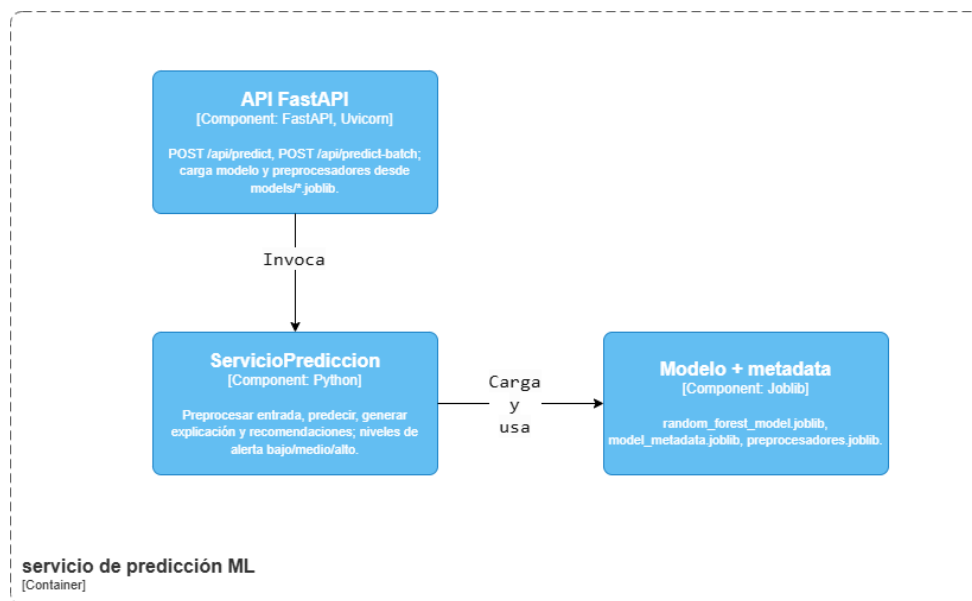


Figura 5.9: Diagrama componente predicción - Fase 1

5.3.6. Modelo de Datos Relacional

El modelo de datos relacional definido para la Fase 1 fue diseñado a partir de los requerimientos funcionales identificados durante el análisis del problema y del proceso de monitoreo deportivo implementado en el caso de estudio. Su propósito principal consistió en centralizar la información de los deportistas, las sesiones de entrenamiento y los resultados generados por el componente analítico, garantizando la trazabilidad de los datos y su disponibilidad para consulta y seguimiento.

La estructura del esquema se encuentra organizada alrededor de la entidad *sesion_entrenamiento*, la cual actúa como núcleo del modelo al relacionar la información del deportista con las diferentes dimensiones de seguimiento deportivo. A partir de esta entidad se vinculan los registros asociados al bienestar emocional, estado fisiológico, historial de lesiones y métricas de rendimiento, permitiendo almacenar de forma estructurada las variables utilizadas durante el proceso de monitoreo.

Adicionalmente, el modelo incorpora las entidades necesarias para la gestión de usuarios y perfiles del sistema, incluyendo entrenadores, deportistas y acudientes, así como las estructuras destinadas al almacenamiento de los resultados generados por el modelo predictivo. Las entidades de predicción, alerta y factores de riesgo permiten conservar la trazabilidad de los análisis realizados y de las notificaciones presentadas a los usuarios dentro de la plataforma.

5.3.7. Resultados de la Implementación de la Fase 1

Las funcionalidades implementadas permitieron validar el flujo completo de información desde la captura de datos en la interfaz de usuario, su procesamiento por parte del backend y la posterior generación de resultados analíticos que son presentados a los distintos actores del sistema. De esta manera, el prototipo constituyó un medio efectivo para obtener retroalimentación de los usuarios participantes en el caso de estudio y para identificar oportunidades de mejora tanto funcionales como arquitectónicas.

A continuación, se presentan algunas de las principales interfaces desarrolladas durante esta fase, las cuales evidencian el estado funcional alcanzado por la plataforma y sirven como base para las mejoras y evolución arquitectónica propuestas en la siguiente iteración del proyecto.

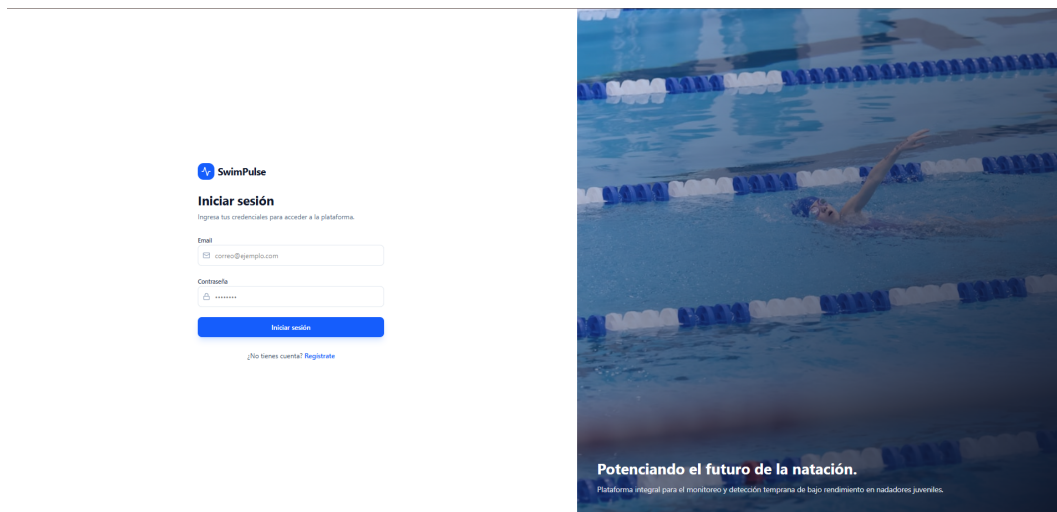


Figura 5.11: Pantalla de autenticación de usuarios

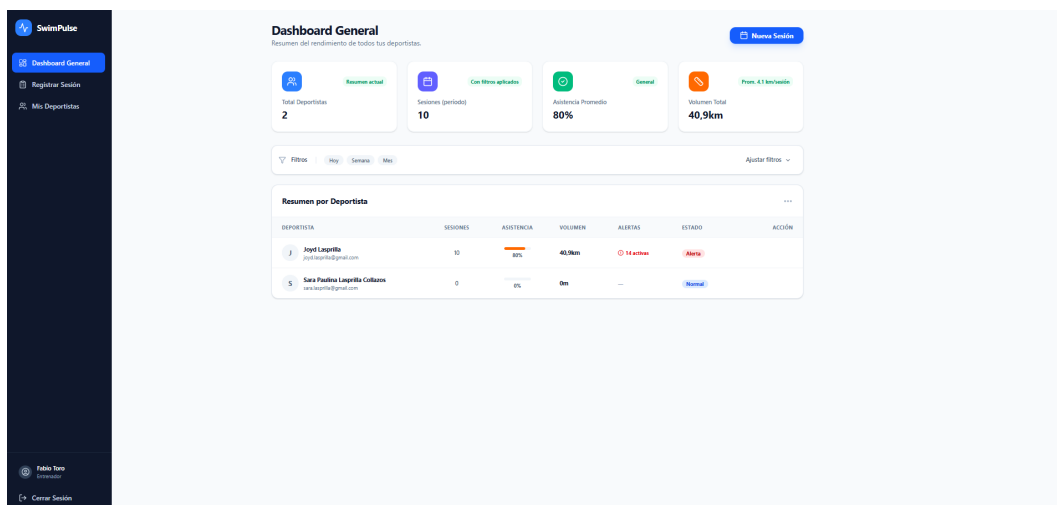


Figura 5.12: Tablero de control para el entrenador

Figura 5.13: Formulario de registro de sesión de entrenamiento



La implementación desarrollada en la Fase 1 permitió validar los procesos funcionales del sistema, obtener retroalimentación de los usuarios involucrados en el caso de estudio e identificar oportunidades de mejora tanto a nivel funcional como arquitectónico. Si bien la solución actual se encuentra estructurada bajo una arquitectura monolítica por capas, la separación de responsabilidades alcanzada entre sus componentes proporciona una base adecuada para evolucionar hacia una arquitectura modular.

A partir de los hallazgos obtenidos durante esta primera iteración, se identificó la necesidad de fortalecer atributos de calidad como la mantenibilidad, la escalabilidad y la capacidad de evolución del sistema. En consecuencia, la Fase 2 plantea una reorganización arquitectónica basada en módulos de negocio claramente delimitados, permitiendo una mayor cohesión funcional, una reducción del acoplamiento entre componentes y una estructura más adecuada para la incorporación de nuevas capacidades y mejoras futuras.

5.4. Fase 2

5.4.1. Trazabilidad entre requerimientos y decisiones arquitectónicas

Las decisiones arquitectónicas adoptadas durante el diseño y evolución de la plataforma no fueron seleccionadas únicamente por criterios tecnológicos, sino como respuesta directa a los requerimientos funcionales y atributos de calidad definidos para el proyecto.

La arquitectura inicial implementada durante la Fase 1 permitió validar los procesos fundamentales del negocio mediante una estructura monolítica por capas. Sin embargo, el crecimiento funcional de la solución, la incorporación del componente analítico y la necesidad de facilitar la evolución futura del sistema evidenciaron la conveniencia de adoptar una organización más alineada con los dominios de negocio.

En consecuencia, durante la Fase 2 se implementó una arquitectura de Monolito Modular basada en principios de Domain-Driven Design (DDD). Esta decisión fue motivada principalmente por los requerimientos no funcionales de modularidad, mantenibilidad y escalabilidad, así como por la necesidad de gestionar de manera independiente los distintos procesos del sistema, tales como la gestión de usuarios, el registro de sesiones de entrenamiento, la analítica predictiva y la visualización de resultados.

De manera complementaria, los requerimientos relacionados con la ejecución de modelos analíticos y generación de alertas tempranas justificaron la incorporación de un servicio de analítica independiente desarrollado en Python. Esta decisión permitió desacoplar la lógica transaccional de la plataforma de los procesos de aprendizaje automático, facilitando la evolución independiente de ambos componentes y aprovechando tecnologías especializadas para ciencia de datos e inteligencia artificial.

Adicionalmente, se evaluó la posibilidad de adoptar una arquitectura basada en microservicios. Sin embargo, esta alternativa no fue seleccionada debido a que los requerimientos actuales del proyecto no justificaban el nivel de complejidad técnica y operativa asociado a dicho estilo arquitectónico. La plataforma presenta un volumen de usuarios moderado, dominios funcionales estrechamente relacionados y una necesidad prioritaria de validación funcional dentro del contexto del caso de estudio. En este escenario, una arquitectura de microservicios habría introducido desafíos adicionales relacionados con la gestión de múltiples despliegues, comunicación distribuida, consistencia de datos, monitoreo y observabilidad, sin aportar beneficios proporcionales a las necesidades actuales del sistema. Por esta razón, se optó por una arquitectura de Monolito Modular basada en Domain-Driven Design (DDD), la cual proporciona los beneficios de separación por dominios, alta cohesión y bajo acoplamiento, manteniendo una complejidad operativa significativamente menor. Asimismo, esta decisión establece una ruta de evolución natural hacia microservicios en escenarios futuros donde el crecimiento funcional, la demanda de escalamiento independiente o la necesidad de despliegues autónomos lo requieran.

La Tabla 5.1 presenta la relación entre los principales requerimientos identificados y las decisiones arquitectónicas implementadas para satisfacerlos.

Tabla 5.1: Trazabilidad entre requerimientos y decisiones arquitectónicas

Requerimientos	Decisión arquitectónica	Justificación
RF01–RF05	Base de datos centralizada y módulo de sesiones de entrenamiento	Permite consolidar toda la información deportiva en una única fuente de datos, garantizando trazabilidad, consulta histórica y consistencia de la información.
RF06–RF09	Módulo de usuarios y autenticación basado en roles	Permite controlar el acceso a funcionalidades y datos según el perfil de entrenador, deportista o acudiente, garantizando seguridad y personalización de la experiencia.
RF10–RF13	Servicio analítico independiente y módulo de analítica	La ejecución de modelos predictivos requiere tecnologías especializadas y ciclos de evolución diferentes al núcleo transaccional del sistema. El desacoplamiento facilita su mantenimiento y escalamiento independiente.
RF14–RF17	Módulo de dashboards y reportes	Permite consolidar métricas, indicadores y resultados analíticos mediante visualizaciones orientadas al apoyo de la toma de decisiones.
RF18–RF21	Integración entre módulo analítico y módulo de alertas	Facilita la generación automática de alertas tempranas a partir de los resultados obtenidos por el modelo predictivo.
RF22–RF24	Arquitectura orientada a múltiples actores	Permite compartir información relevante entre entrenadores, deportistas y acudientes de acuerdo con sus permisos y necesidades de seguimiento.
Modularidad, Mantenibilidad y Escalabilidad	Monolito Modular basado en DDD	La separación por dominios funcionales reduce el acoplamiento entre componentes, incrementa la cohesión interna y facilita la evolución progresiva de la plataforma.
Complejidad operativa controlada y evolución progresiva	Monolito Modular en lugar de Microservicios	Permite obtener modularidad, mantenibilidad y separación de dominios sin introducir los costos de infraestructura, despliegue distribuido y gestión operacional propios de una arquitectura de microservicios.

Los resultados de la trazabilidad evidencian como la arquitectura propuesta responde directamente a los requerimientos funcionales y no funcionales definidos para el proyecto. Particularmente, la adopción de una arquitectura de Monolito Modular basada en DDD permitió alinear la estructura técnica con los procesos del negocio, mientras que la incorporación de un servicio analítico independiente facilitó la integración de capacidades de inteligencia artificial sin afectar la estabilidad del núcleo transaccional de la plataforma.

Esta alineación entre requerimientos y decisiones arquitectónicas proporciona una base sólida para la evolución futura del sistema y demuestra que las decisiones de diseño adoptadas no obedecen únicamente a criterios tecnológicos, sino a la necesidad de satisfacer los objetivos funcionales y atributos de calidad establecidos para la solución.

5.4.2. Enfoque arquitectónico

A partir de los resultados obtenidos durante la primera fase y de la retroalimentación recopilada en el caso de estudio, se identificó la necesidad de evolucionar la arquitectura del sistema hacia una estructura que facilitara su crecimiento, mantenimiento y evolución funcional. Si bien la arquitectura monolítica por capas permitió validar los procesos principales del negocio y acelerar la construcción del prototipo inicial, el incremento de responsabilidades dentro de un único esquema organizacional evidenció la necesidad de fortalecer la cohesión funcional y reducir el acoplamiento entre componentes.

En respuesta a esta necesidad, la segunda fase adopta una arquitectura de tipo Monolito Modular basada en principios de Domain-Driven Design (DDD), donde la organización del sistema deja de centrarse en capas técnicas globales y pasa a estructurarse alrededor de dominios funcionales claramente delimitados. Este enfoque permite encapsular reglas de negocio, casos de uso, mecanismos de persistencia y contratos de exposición dentro de cada módulo, favoreciendo una mayor independencia funcional y facilitando la incorporación de nuevas capacidades sin afectar el comportamiento de los demás componentes.

La arquitectura conserva las ventajas operativas de un despliegue monolítico, pero introduce una organización interna orientada al negocio que fortalece atributos de calidad como mantenibilidad, modificabilidad, escalabilidad funcional y capacidad de evolución. Asimismo, establece una base arquitectónica que permite una futura transición hacia arquitecturas distribuidas si el crecimiento de la plataforma así lo requiere.

5.4.3. Arquitectura Monolito Modular basada en DDD

La arquitectura del backend fue reorganizada siguiendo los principios de Domain-Driven Design (DDD), agrupando los componentes según las capacidades funcionales del negocio en lugar de hacerlo únicamente por responsabilidad técnica. Bajo este enfoque, cada módulo constituye una unidad funcional cohesionada que encapsula sus propias interfaces de entrada, lógica de aplicación, modelo de dominio y mecanismos de persistencia.

Los principales módulos identificados dentro de la plataforma son:

- **Módulo de Usuarios y Autenticación:** responsable de la gestión de usuarios, autenticación, autorización y administración de perfiles según los diferentes roles del sistema.
- **Módulo de Sesiones de Entrenamiento:** concentra la gestión de sesiones deportivas, el registro de métricas de entrenamiento, variables fisiológicas y factores asociados al desempeño del deportista.
- **Módulo de Analítica y Predicción:** encapsula la integración con el servicio externo de inteligencia artificial, gestionando las solicitudes de predicción, el procesamiento de resultados y la persistencia de los indicadores generados.
- **Módulo de Dashboard y Reportes:** consolida la información requerida para la visualización de métricas, tendencias, alertas e indicadores de rendimiento dirigidos a entrenadores, deportistas y acudientes.
- **Módulo de Configuración y Listas de Referencia:** administra la información parametrizable utilizada por los distintos procesos de negocio de la plataforma.

Cada módulo se estructura internamente mediante las capas de API, Aplicación, Dominio e Infraestructura, promoviendo una separación clara de responsabilidades y limitando las dependencias entre contextos funcionales. Esta organización permite que las funcionalidades evolucionen de manera independiente, simplifica las actividades de mantenimiento y proporciona una arquitectura alineada con los objetivos de modularidad definidos para esta segunda fase.

Obteniendo como resultado el diagrama de componentes del backend, donde se visualiza la organización modular de la solución y las relaciones existentes entre los distintos módulos de negocio que conforman la plataforma.

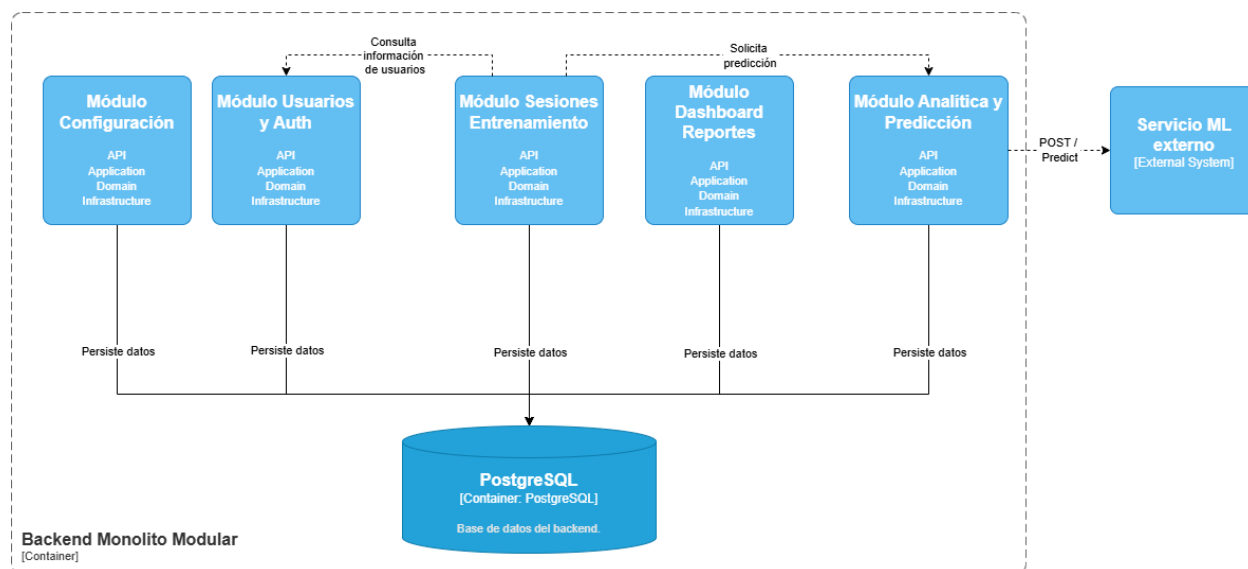


Figura 5.15: Diagrama de componentes del backend bajo arquitectura Monolito Modular basada en DDD

5.4.4. Integración del servicio de analítica y predicción

La plataforma integra un servicio externo de analítica y predicción desarrollado en Python con FastAPI, consumido por el backend mediante una invocación HTTP en el flujo de registro de sesión de entrenamiento. El servicio define dos *endpoints* públicos: `GET /health`, utilizado para verificar que el modelo y sus artefactos auxiliares (escalas, codificadores y metadatos) se encuentran correctamente cargados antes de delegar predicciones, y `POST /predict`, responsable de la inferencia. Cuando se registra una sesión con asistencia confirmada, el backend envía a `POST /predict` un objeto que combina los identificadores del deportista, estilo, tipo de sesión y estado de ánimo, las variables de bienestar y rendimiento de la sesión actual (RPE, horas de sueño, nivel de estrés, motivación, presencia de dolor, volumen total, tiempo real y tiempo planificado) y el historial de los tres registros más recientes de tiempo, RPE y horas de sueño, a partir del cual el servicio construye internamente las variables temporales que el modelo requiere.

Internamente, el *endpoint* consolida en una sola operación los resultados de las dos fases analíticas del proyecto. Primero arma el conjunto de características de entrada, las escala, ejecuta el modelo MLP entrenado en la Fase 2 y revierte la transformación de escala para obtener el tiempo real esperado de la siguiente sesión en segundos. A continuación evalúa los siete criterios de riesgo definidos en la Fase 1 sobre la misma sesión y produce una clasificación del estado del deportista compuesta por una bandera de bajo rendimiento, una puntuación entera entre cero y siete, una probabilidad normalizada y un nivel categórico de riesgo en cuatro escalones (bajo, medio, alto y muy alto). La respuesta incluye además el detalle explicable de cada uno de los siete factores,

indicando si se encuentra activo en la sesión, su valor actual, el umbral aplicado y una descripción interpretable. Con base en esta información, el backend persiste la predicción y genera las entidades de alerta y factores de riesgo asociadas, que posteriormente se exponen al usuario a través de los tableros de control y las vistas de detalle de sesión.

5.4.5. Modelo de Datos Relacional

Durante la Fase 2, el modelo de datos evolucionó con el objetivo de mejorar la representación del contexto deportivo del atleta, optimizar el almacenamiento de la información utilizada por el componente analítico y fortalecer la trazabilidad de los procesos implementados en la plataforma. Aunque se mantiene la entidad *sesion_entrenamiento* como eje central del esquema, se realizaron ajustes estructurales orientados a soportar los nuevos requerimientos funcionales y analíticos identificados durante la validación del prototipo.

Entre las principales modificaciones se encuentra la incorporación de entidades de referencia para la gestión de estilos de nado, tipos de sesión, estados de ánimo y niveles de riesgo, así como la inclusión de la entidad *historial_sesion*, encargada de almacenar información histórica relevante para la ejecución del modelo predictivo. Asimismo, se fortaleció la gestión de usuarios mediante la separación de roles en entidades especializadas, proporcionando una estructura más flexible y alineada con los principios de modularidad adoptados en esta fase.

La integración entre las entidades de sesión de entrenamiento, historial deportivo, predicción, alertas y factores de riesgo permite consolidar en un único modelo la información operativa y analítica necesaria para el monitoreo del rendimiento deportivo y la generación de alertas tempranas. Esta evolución del esquema contribuye a mejorar la consistencia de los datos, facilitar futuras extensiones funcionales y soportar de manera más eficiente los procesos de inteligencia artificial incorporados en la plataforma.

A continuación se presenta el Modelo relacional correspondiente a la Fase 2, evidenciando la evolución del esquema respecto a la versión implementada en la fase anterior.

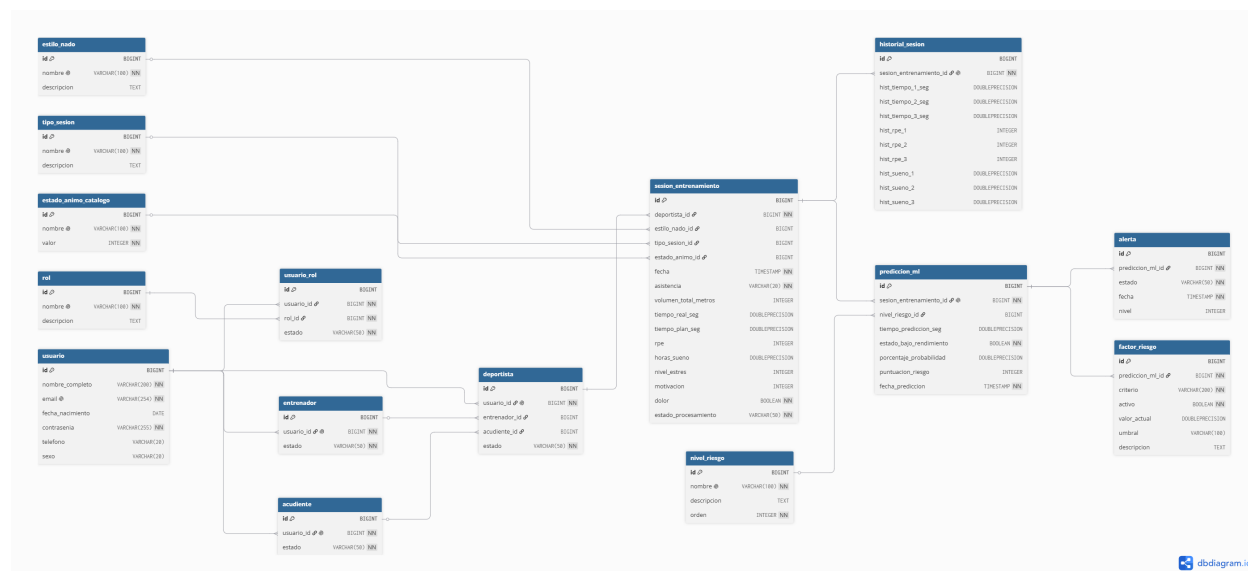


Figura 5.16: Modelo relacional de la plataforma en la Fase 2

La evolución del modelo de datos refleja la madurez alcanzada por la solución durante la segunda fase de desarrollo. Las modificaciones realizadas no solo permitieron soportar los nuevos requerimientos funcionales y analíticos del sistema, sino también mejorar atributos de calidad como la mantenibilidad, la escalabilidad y la capacidad de evolución futura de la plataforma.

5.4.6. Resultados de la Implementación de la Fase 2

Como resultado de la segunda fase de desarrollo, la plataforma evolucionó tanto a nivel funcional como arquitectónico. A partir de la retroalimentación obtenida durante la validación del prototipo inicial, se incorporaron nuevas capacidades orientadas a fortalecer la centralización de datos, mejorar los mecanismos de análisis del rendimiento deportivo y consolidar una arquitectura interna alineada con los principios de modularidad, mantenibilidad y escalabilidad definidos para el proyecto.

Desde la perspectiva funcional, la plataforma amplió la cobertura de los procesos de gestión, consulta y seguimiento del desempeño deportivo. Se fortaleció la integración con el componente analítico mediante la incorporación del modelo predictivo desarrollado en la investigación, permitiendo generar estimaciones de desempeño, identificar condiciones asociadas al bajo rendimiento y emitir alertas tempranas que apoyan la toma de decisiones por parte de entrenadores, deportistas y académicos. Asimismo, se mejoró la presentación de la información mediante tableros de control e interfaces orientadas a facilitar la interpretación de indicadores y resultados analíticos.

Desde la perspectiva arquitectónica, el backend fue reorganizado bajo un enfoque de Monolito Modular basado en principios de Domain-Driven Design (DDD), estructurando la solución en módulos funcionales independientes para la gestión de usuarios, sesiones de entrenamiento, analíti-

ca y predicción, dashboards y configuración. Esta reorganización permitió reducir el acoplamiento entre componentes, mejorar la cohesión funcional y establecer una base arquitectónica preparada para la evolución futura del sistema.

A continuación, se presentan algunas de las principales funcionalidades e interfaces desarrolladas durante esta fase, las cuales evidencian la evolución alcanzada por la plataforma y los resultados obtenidos durante la implementación del prototipo.

El módulo de registro de sesión (Figura 5.17) consolidó la captura estructurada de las variables de carga, percepción y rendimiento, integrándose con el componente analítico para alimentar de forma directa los modelos predictivos.

SwimPulse

Dashboard

Registrar sesión

Deportistas

Alertas

Fabio Toro
Entrenador

Cerrar sesión

Registrar sesión

Captura los datos de la sesión de entrenamiento

Identificación de la sesión

Deportista
Seleccionar...

Fecha y hora de la sesión
13/06/2026 03:11 p. m.

Estilo de nado
Seleccionar...

Tipo de sesión
Seleccionar...

Asistencia
☒ Presente ☐ Ausente

Datos de rendimiento

Tiempo planificado 100m (min:seg)
1:08
Formato min:seg (ej. 1:08 = 68 s)

Tiempo real 100m (min:seg)
1:10
Formato min:seg (ej. 1:08 = 68 s)

Volumen total (metros)
2000

Carga interna y bienestar

RPE: 5 — Duro

Horas de sueño
8

Nivel de estrés: 3 — Moderado

Motivación: 7 — Alta

Presencia de dolor
No

Estado de ánimo
Seleccionar...

Cancelar Registrar sesión

Figura 5.17: Registro de sesión de entrenamiento con integración al módulo analítico.

El dashboard del entrenador (Figura 5.18) centralizó los indicadores consolidados de rendimiento del plantel, facilitando la interpretación del desempeño y el seguimiento de cada deportista a partir de las estimaciones del modelo.

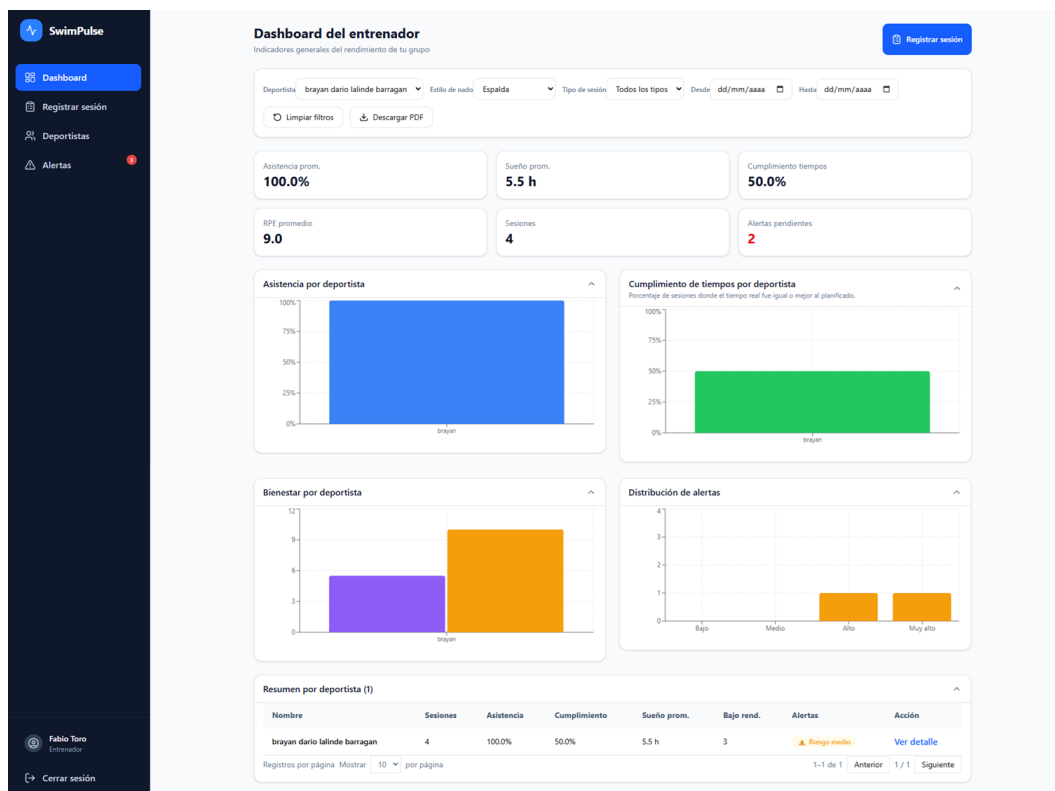


Figura 5.18: Dashboard del entrenador con indicadores consolidados de rendimiento.

Finalmente, el módulo de alertas (Figura 5.19) materializó la detección temprana del bajo rendimiento, presentando de forma visible las condiciones de riesgo identificadas por el modelo para apoyar la toma de decisiones del cuerpo técnico.

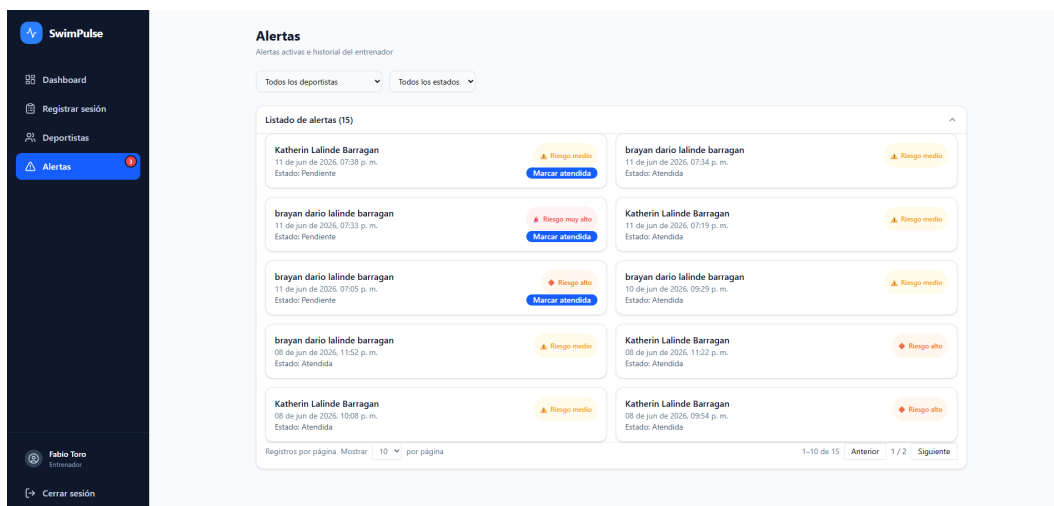


Figura 5.19: Módulo de alertas tempranas con las condiciones de riesgo identificadas por el modelo predictivo.

5.4.7. Proyección a Escalamiento de la Arquitectura

La arquitectura definida en la Fase 2 establece una base sólida para el crecimiento progresivo de la plataforma, gracias a la organización modular del backend y a la separación de responsabilidades entre los distintos componentes del sistema. La adopción de un enfoque de Monolito Modular basado en Domain-Driven Design (DDD) permite que cada módulo evolucione de forma controlada, incorporando nuevas funcionalidades o modificaciones sin generar impactos significativos sobre el resto de la solución.

La delimitación de contextos funcionales y el desacoplamiento existente entre la plataforma web, el backend, la base de datos y el servicio de analítica crean condiciones favorables para incrementar la capacidad operativa del sistema a medida que aumente el número de usuarios, el volumen de información procesada o la complejidad de los modelos predictivos utilizados.

Adicionalmente, la integración del componente analítico mediante interfaces HTTP permite que este evolucione de manera independiente, facilitando la incorporación de nuevos modelos de inteligencia artificial, algoritmos de predicción o servicios especializados sin requerir modificaciones estructurales sobre la plataforma principal.

En escenarios de crecimiento futuro, la arquitectura modular implementada proporciona una ruta natural de evolución hacia esquemas distribuidos, donde los módulos con mayor demanda o autonomía funcional podrían desplegarse como servicios independientes. De esta manera, la solución desarrollada no solo satisface los requerimientos actuales del caso de estudio, sino que establece las bases arquitectónicas necesarias para soportar procesos de expansión, nuevas capacidades funcionales y mayores exigencias operativas en etapas posteriores del proyecto.

5.4.8. Evaluación de la Evolución Arquitectónica

La evolución arquitectónica realizada entre la Fase 1 y la Fase 2 tuvo como objetivo principal mejorar la organización interna del sistema, incrementar la cohesión funcional de sus componentes y reducir el acoplamiento entre las distintas responsabilidades de negocio. Aunque ambas fases conservan el mismo contexto funcional y tecnológico, la segunda iteración introduce una reorganización estructural basada en principios de Domain-Driven Design (DDD), permitiendo una mejor alineación entre la arquitectura del software y los dominios del negocio.

La evaluación arquitectónica se realizó mediante una revisión comparativa de la estructura del software en ambas fases, utilizando como referencia los atributos de calidad definidos en los requerimientos no funcionales del proyecto. Para cada atributo se analizaron aspectos como la organización de dependencias, encapsulamiento de responsabilidades, impacto esperado de cambios funcionales y capacidad de extensión de los módulos existentes. La Tabla 5.2 resume las principales diferencias identificadas entre ambas fases.

Tabla 5.2: Comparación arquitectónica entre la Fase 1 y la Fase 2

Aspecto	Fase 1	Fase 2
Estilo arquitectónico	Monolito por capas	Monolito Modular basado en DDD
Organización interna	Agrupación por responsabilidades técnicas	Agrupación por dominios funcionales de negocio
Cohesión funcional	Media	Alta
Acoplamiento entre componentes	Medio	Bajo
Mantenibilidad	Modificaciones distribuidas entre capas compartidas	Cambios localizados dentro de cada módulo
Escalabilidad funcional	Limitada por la organización global del sistema	Facilitada mediante módulos independientes
Integración del componente analítico	Servicio externo desacoplado	Servicio externo desacoplado con responsabilidades encapsuladas en un módulo específico
Capacidad de evolución	Requiere mayor conocimiento transversal del sistema	Permite evolución controlada por contexto funcional
Preparación para crecimiento futuro	Parcial	Alta

Los resultados obtenidos evidencian que la arquitectura propuesta en la Fase 2 mejora los atributos de calidad definidos para el proyecto sin afectar las funcionalidades previamente implementadas. La reorganización modular permitió delimitar claramente los contextos funcionales del sistema, sim-

plificar las actividades de mantenimiento y establecer una base arquitectónica más adecuada para la incorporación futura de nuevas capacidades analíticas, mecanismos de comunicación y procesos de escalamiento.

5.5. Resumen del capítulo

El capítulo presentó el diseño y desarrollo del prototipo en dos fases evolutivas. Se definieron los requerimientos funcionales y no funcionales del sistema, y se justificó el stack tecnológico seleccionado: React/TypeScript en el frontend, Spring Boot/Java en el backend, PostgreSQL para la persistencia y Python/FastAPI para el componente analítico. En la Fase 1 se construyó un prototipo funcional con arquitectura monolítica por capas, que permitió validar los procesos fundamentales del sistema e identificar oportunidades de mejora funcionales y arquitectónicas. En la Fase 2, a partir de esos hallazgos, el backend fue reorganizado bajo un enfoque de Monolito Modular basado en Domain-Driven Design (DDD), estructurando la solución en dominios funcionales claramente delimitados, lo que incrementó la cohesión interna, redujo el acoplamiento y fortaleció atributos de calidad como mantenibilidad, modificabilidad y escalabilidad funcional. Se justificó este enfoque frente a los microservicios, dado que el contexto del proyecto no justificaba la complejidad operativa de un sistema distribuido. Se estableció trazabilidad explícita entre requerimientos y decisiones arquitectónicas, demostrando que el diseño responde a objetivos funcionales y de calidad, no solo a criterios tecnológicos. La evaluación comparativa entre fases evidenció una mejora significativa en la organización interna y en la preparación del sistema para su crecimiento futuro.

Validación del prototipo en el caso de estudio Fabio Toro Academy

La validación constituyó la etapa de comprobación del proyecto y el mecanismo mediante el cual se contrastaron las decisiones de diseño, desarrollo y modelado con las necesidades reales de los actores del proceso formativo. Su propósito no se limitó a comprobar la operación del software, sino a determinar, en qué medida el trabajo realizado satisface el objetivo general, los objetivos específicos y las preguntas de investigación planteadas. Para ello, la validación se concibió como un proceso iterativo: (Fase 1) validó el prototipo inicial en condiciones reales de uso y reveló oportunidades de mejora, y (Fase 2) verificó que dichas mejoras, tanto en el componente analítico como en la arquitectura, respondieran de manera efectiva a las falencias detectadas.

6.1. Diseño de la validación

La validación se realizó en Fabio Toro Academy, de acuerdo al alcance definido para el proyecto y se articuló en cuatro dimensiones complementarias, cada una orientada a un aspecto distinto del trabajo realizado y sustentada, como resume la Tabla [6.1](#).

Tabla 6.1: Dimensiones del diseño de validación, evidencia e instrumentos.

Dimensión	Qué se validó y con qué evidencia	Participante
Validación funcional	Operación de los flujos principales del sistema (registro de sesiones, centralización del historial, generación de alertas, tableros y reportes), mediante uso guiado del prototipo en las instalaciones de la academia.	Entrenador jefe
Evaluación técnica del componente analítico	Adecuación del enfoque analítico al objetivo y desempeño del modelo, mediante juicio de experto y revisión de las métricas de evaluación (Capítulo 4 y Anexo B).	Experto en inteligencia artificial
Validación del insumo de datos	Correspondencia del conjunto de datos sintético contextualizado con las condiciones operativas del caso, mediante contraste con la programación real del semestre 2026-1 y historial deportivo de los 14 atletas de la academia	Equipo de desarrollo y entrenador jefe
Utilidad percibida y apoyo a la decisión	Aporte del prototipo al seguimiento, al alertamiento y a la comunicación entre actores, mediante comparación cualitativa del proceso antes y después de su incorporación.	Entrenador jefe

Nota. La validación corresponde a un caso de estudio único con evaluación cualitativa por juicio de experto; sus alcances y límites se discuten en la Sección 6.5.

Como referencia de contraste se adoptaron los criterios de éxito definidos para el proyecto: **(C1)** centralización efectiva de la información de entrenamiento; **(C2)** integración operativa del modelo analítico dentro del flujo del sistema; **(C3)** generación automática de alertas tempranas de bajo rendimiento; **(C4)** apoyo a la toma de decisiones y a la comunicación entre entrenador, deportistas y acudientes; y **(C5)** viabilidad técnica y sostenibilidad arquitectónica del prototipo. Estos criterios guían la verificación del cumplimiento de objetivos presentada en la Sección 6.4.

6.2. Validación de la Fase 1: prototipo inicial

La validación de la Fase 1 se llevó a cabo en las instalaciones de Fabio Toro Academy e integró dos dimensiones: la validación funcional de los flujos del sistema con el entrenador jefe y la evaluación técnica del componente analítico con un experto en inteligencia artificial.

6.2.1. Validación funcional con el entrenador jefe

La validación funcional confirmó la correcta operación de los flujos principales del prototipo. El entrenador jefe valoró tres aspectos como aportes relevantes para su práctica diaria: el mecanismo de alertamiento automático tras la inferencia, que expone la predicción binaria de bajo rendimiento, el nivel categórico de riesgo y el detalle explicable de los siete factores evaluados; la generación

de reportes en PDF con filtrado por nadador, período y variable, como canal estructurado de comunicación con deportistas, acudientes y directivos; y la centralización de la información de cada nadador en un espacio accesible desde distintas vistas, en contraste con la gestión dispersa preexistente. El entrenador validó, adicionalmente, la pertinencia de los siete factores de riesgo del modelo y solicitó explícitamente su conservación para el alertamiento en fases subsiguientes.

Los flujos ejercitados durante esta validación corresponden a los grupos de requerimientos funcionales definidos en el capítulo de diseño del prototipo: centralización de datos (RF01–RF05), procesamiento y resultados del modelo analítico (RF10–RF13), detección y emisión de alertas (RF18–RF21), visualización mediante tableros (RF14–RF17) y apoyo a la comunicación entre actores (RF22–RF24). La validación funcional confirmó la operación de estos flujos en el caso de estudio; no obstante, no se aplicó un protocolo formal de verificación requerimiento por requerimiento ni instrumentos cuantitativos de usabilidad, aspecto que se retoma en las limitaciones y en los trabajos futuros.

6.2.2. Evaluación técnica del componente analítico

La evaluación técnica identificó que el modelo implementado en la Fase 1 resolvía un problema de clasificación binaria y no uno de predicción en sentido estricto, lo que constituía un posible incumplimiento del objetivo analítico, que contemplaba la estimación continua del desempeño individual futuro del nadador. A partir de este hallazgo, el experto recomendó incorporar un Perceptrón Multicapa (MLP) orientado a la estimación del tiempo esperado por nadador en la sesión siguiente. Esta reorientación tuvo implicaciones directas sobre la arquitectura del sistema: la incorporación del MLP exigía gestionar el historial reciente de sesiones por nadador y diferenciar el estado de procesamiento de cada registro, necesidades que la estructura monolítica en capas de la Fase 1 no contemplaba.

Ante la baja adherencia de los nadadores al formulario de captura se recomendó, además, la construcción de un conjunto de datos sintético contextualizado para sustentar el entrenamiento del modelo en la Fase 2.

6.2.3. Síntesis de hallazgos y decisiones de rediseño

La Fase 1 cumplió su función validadora con un doble resultado. Por un lado, confirmó el valor operativo del prototipo en los flujos centrales del caso de estudio; por otro, evidenció dos brechas que orientaron la segunda iteración: la necesidad de transformar la salida analítica de una clasificación binaria a una predicción continua del rendimiento, y la conveniencia de disponer de un insumo de datos suficiente para entrenar y evaluar dicho modelo. Estas dos decisiones constituyen el puente metodológico entre ambas fases y demuestran el carácter formativo de la validación, que no se limitó a constatar resultados, sino que reorientó el desarrollo.

6.3. Validación de la Fase 2

La validación de la Fase 2 verificó que las mejoras incorporadas en el componente analítico, en el modelo de datos y en la arquitectura del backend dieran respuesta integral a los hallazgos de la iteración anterior.

6.3.1. Validación del conjunto de datos sintético contextualizado

El conjunto de datos sintético fue validado en cuanto a su correspondencia con el caso de uso: sus supuestos de generación, distribuciones de variables y criterios de etiquetado se elaboraron a partir de la programación real del semestre 2026-1 de Fabio Toro Academy e historial deportivo de los 14 atletas de la academia, lo que otorgó al conjunto una base contextual coherente con las condiciones operativas del sistema. El detalle del pipeline de generación y depuración se documenta en el Anexo A. Esta validación confirma que el insumo de datos es metodológicamente trazable y pertinente para el caso; no sustituye, sin embargo, a los datos reales, por lo que las métricas derivadas de él se interpretan como evidencia de *viabilidad técnica* y no como caracterización empírica definitiva del rendimiento (Carvalho et al., 2024).

6.3.2. Verificación del cumplimiento del objetivo analítico

La incorporación del MLP permitió cerrar la brecha detectada en la Fase 1. El modelo fue evaluado en su capacidad de estimar el tiempo real de 100 m esperado en la sesión siguiente, obteniendo en el conjunto de prueba un MAE de 3,688 segundos, un RMSE de 4,547 segundos y un R^2 de 0,778, cifras ya analizadas en el Capítulo 4 (Tabla 4.3) y sustentadas gráficamente en el Anexo B (Figuras B.9, B.10 y B.11). Estas métricas reflejan una capacidad predictiva operativamente útil para una estimación de corto plazo en un contexto formativo, reconociendo que operan sobre datos sintéticos y, por tanto, describen la viabilidad técnica del enfoque y no su precisión definitiva en condiciones reales (Carvalho et al., 2024). Se identificó un desfase sistemático de aproximadamente +2,55 segundos, asociado a la alta correlación ($r = 0,97$) entre el tiempo planificado y la variable objetivo; este sesgo es conservador para un sistema de alerta y se corregirá a medida que se incorporen datos empíricos al sistema.

Desde la perspectiva de la validación, el resultado relevante es que la transición de una salida binaria a una predicción continua del tiempo esperado satisface la formulación del objetivo analítico que la Fase 1 no alcanzaba a cubrir en sentido estricto. La validación confirmó, adicionalmente, que la integración del MLP con los siete factores de riesgo conservados de la Fase 1 permite concretar el alertamiento temprano de forma integral: el sistema combina la estimación del rendimiento futuro con la evaluación del estado actual del deportista, generando alertas fundamentadas tanto en la magnitud del tiempo estimado como en los indicadores de bienestar y carga activos en la sesión.

6.3.3. Validación de la evolución arquitectónica

La Fase 2 materializó la proyección modular anticipada en el diseño de la Fase 1. La arquitectura centralizada en capas evolucionó hacia una organización modular del backend con responsabilidades claramente diferenciadas. Aunque el backend se mantiene como un único artefacto de compilación, la separación interna por dominios funcionales establece límites lógicos que pueden transformarse en límites físicos en una etapa posterior sin reescrituras completas, siguiendo una estrategia de *arquitectura modular evolutiva*.

El cambio arquitectónico más crítico, verificado durante la validación, fue la incorporación de una capa dedicada de integración analítica con un mecanismo de contingencia explícito: ante un fallo del servicio externo, el sistema persiste una predicción mínima que preserva el registro de la sesión sin interrumpir al usuario, mientras la columna de estado de procesamiento introducida en el modelo de datos permite diferenciar sesiones procesadas, pendientes o fallidas, habilitando su reprocesamiento posterior. Esta decisión resolvió directamente el riesgo operativo identificado en la Fase 1 y garantiza la continuidad del registro de entrenamiento, requisito no negociable para el cuerpo técnico. El modelo de datos relacional fue ajustado, además, para incorporar el historial de los tres registros más recientes de tiempo, RPE y horas de sueño por nadador, insumo requerido por el MLP para la construcción de variables temporales. La validación confirmó que la nueva organización modular opera correctamente y que la trazabilidad del ciclo de vida de cada sesión garantiza la coherencia de los datos que alimentan los tableros de control y los reportes del cuerpo técnico.

6.3.4. Evaluación del apoyo a la toma de decisiones y la gestión formativa

La validación permitió analizar no solo el comportamiento técnico del sistema, sino también su aporte a la toma de decisiones dentro del proceso formativo de la academia. Para ello se comparó el escenario previo a la implementación del prototipo con el escenario posterior a su utilización durante las sesiones de validación funcional.

Antes de la implementación, la información relacionada con asistencia, bienestar, rendimiento y carga de entrenamiento se encontraba distribuida en diferentes fuentes y dependía principalmente de registros manuales y de la interpretación individual del entrenador. Esta situación dificultaba la consolidación histórica de la información, la identificación temprana de patrones de deterioro del rendimiento y la comunicación estructurada con deportistas y acudientes. Tras la incorporación del prototipo, la información pasó a centralizarse en una única plataforma, lo que permitió consultar el historial de sesiones, visualizar indicadores consolidados mediante tableros y acceder a reportes generados automáticamente. De acuerdo con la validación realizada con el entrenador jefe, esta integración facilitó la interpretación conjunta de variables que antes debían analizarse por separado y redujo el esfuerzo requerido para el seguimiento individual de cada nadador. La Tabla 6.2 sintetiza esta comparación cualitativa.

Tabla 6.2: Comparación cualitativa del proceso de seguimiento antes y después del prototipo.

Aspecto	Situación previa	Situación con el prototipo
Almacenamiento de información	Registros dispersos en múltiples fuentes	Información centralizada en una única plataforma
Seguimiento del rendimiento	Dependiente de revisión manual	Indicadores consolidados y tableros interactivos
Detección de riesgo	Basada en observación subjetiva	Alertamiento automático basado en variables registradas
Proyección de rendimiento	No disponible	Estimación del tiempo esperado mediante modelo predictivo
Comunicación con acudientes	Reportes elaborados manualmente	Generación de reportes en PDF
Consulta histórica	Limitada y fragmentada	Historial consolidado por nadador

Desde la perspectiva del monitoreo deportivo, el prototipo incorporó dos mecanismos complementarios de apoyo a la decisión: el alertamiento basado en factores de riesgo y la estimación del rendimiento futuro mediante el modelo predictivo. Esta combinación permitió disponer simultáneamente de información descriptiva sobre el estado actual del deportista y de una proyección sobre su posible comportamiento en la siguiente sesión, ampliando el contexto disponible para la planificación de cargas y el seguimiento formativo. Asimismo, la generación automática de reportes y la disponibilidad permanente de la información favorecieron la comunicación entre entrenador, deportistas y acudientes, transformando datos previamente dispersos en información estructurada e interpretable. Si bien el tiempo disponible para el desarrollo del trabajo no permitió una evaluación longitudinal que midiera cambios en indicadores deportivos a lo largo de varias temporadas, la validación funcional confirmó la utilidad percibida del sistema y su capacidad para proporcionar información relevante, oportuna y contextualizada para la gestión cotidiana del rendimiento.

6.4. Verificación del cumplimiento de los objetivos

La evidencia recogida a lo largo de las dos iteraciones permite verificar el cumplimiento del objetivo general y de los cuatro objetivos específicos. La Tabla 6.3 articula, para cada objetivo, el entregable que lo materializa, la evidencia de validación que lo respalda y el estado alcanzado.

Tabla 6.3: Verificación del cumplimiento de los objetivos del proyecto.

Objetivo	Entregable / resultado	Evidencia de validación	Estado
OE1. Identificar y caracterizar los datos	Caracterización de variables de carga, rendimiento, bienestar y contexto; pipeline de preprocesamiento; conjunto de datos sintético contextualizado (Capítulo 4, Anexo A).	Validación de la pertinencia de las variables con el cuerpo técnico y del conjunto sintético frente a la programación real 2026-1.	Cumplido (con matiz)
OE2. Desarrollar el modelo analítico y predictivo	Clasificación de bajo rendimiento (Fase 1) y predicción continua del tiempo futuro mediante MLP (Fase 2), con sus métricas de evaluación (Capítulo 4, Anexo B).	Evaluación técnica por experto que reorientó el enfoque, y verificación de que el MLP cubre la predicción continua exigida por el objetivo (Sección 6.3.2).	Cumplido (con matiz)
OE3. Desarrollar el prototipo web modular	Prototipo funcional (frontend, backend, base de datos y servicio analítico) con arquitectura que evoluciona a Monolito Modular basado en DDD y mecanismo de contingencia.	Validación funcional de los flujos y verificación de la evolución arquitectónica (Sección 6.3.3).	Cumplido
OE4. Validar la plataforma en el caso de estudio	Validación funcional, técnica y de utilidad percibida en Fabio Toro Academy, con comparación del proceso antes y después.	El presente capítulo en su totalidad.	Cumplido
Objetivo general	Prototipo de plataforma web modular que centraliza datos, integra un modelo analítico, genera alertas tempranas y apoya la toma de decisiones, aplicado al caso de estudio.	Cumplimiento conjunto de OE1–OE4 y de los criterios de éxito C1–C5.	Cumplido

Nota. “Cumplido (con matiz)” señala objetivos alcanzados cuya evidencia descansa parcialmente en datos sintéticos o presenta fuga de información reconocida; los matices se detallan en la Sección 6.5.

En términos de los criterios de éxito definidos a priori, la validación evidenció que el prototipo centraliza la información antes dispersa (C1), integra el modelo analítico dentro del flujo de registro de sesiones (C2), genera alertas tempranas combinando estado actual y predicción de rendimiento (C3), y aporta mecanismos concretos de apoyo a la decisión y a la comunicación entre actores (C4), sobre una arquitectura modular técnicamente viable y preparada para evolucionar (C5). En conjunto, esta evidencia permite responder afirmativamente la pregunta de investigación general: es posible materializar un prototipo de plataforma web modular que centralice y analice los datos de rendimiento para apoyar la detección temprana del bajo rendimiento y la comunicación entre los actores de una academia de formación, según se constató en el caso de estudio. El alcance de esta respuesta queda delimitado por las consideraciones de la sección siguiente.

6.5. Limitaciones y amenazas a la validez

La interpretación de los resultados de validación exige reconocer sus límites, que orientan además los trabajos futuros. En primer lugar, la validación se efectuó sobre un *caso de estudio único*, por lo que sus conclusiones no son directamente generalizables a otras academias sin replicación. En segundo lugar, la evidencia cualitativa se sustenta en el juicio de un entrenador jefe (validación funcional y de utilidad) y de un experto en inteligencia artificial (evaluación técnica); no se aplicaron instrumentos cuantitativos estructurados de usabilidad ni un protocolo de verificación formal requerimiento por requerimiento. En tercer lugar, las métricas predictivas de la Fase 2 se obtuvieron sobre un conjunto de datos *sintético*, cuya señal fue parcialmente construida de forma controlada; describen, por tanto, la viabilidad técnica del enfoque y no su precisión definitiva en condiciones reales. A ello se suma la fuga de información reconocida en la Fase 1 y la alta correlación entre el tiempo planificado y la variable objetivo, factores que pueden inflar las métricas y reducir la visibilidad del aporte de las variables de bienestar y carga. Finalmente, el alcance temporal del proyecto impidió una evaluación longitudinal de impactos deportivos.

Estas limitaciones no invalidan el cumplimiento de los objetivos verificado en la sección anterior, pero sí precisan su alcance: el aporte central del trabajo es una solución de ingeniería de software que articula gestión de datos, aprendizaje automático, arquitectura modular y validación en un entorno real, cuya consolidación empírica dependerá de la recalibración con datos reales y de una validación más amplia, según se detalla en las conclusiones y los trabajos futuros.

6.6. Resumen del capítulo

El capítulo documentó la validación del prototipo en Fabio Toro Academy bajo un diseño explícito de cuatro dimensiones ejecutado de forma iterativa en dos fases. En la Fase 1 se confirmó el valor operativo del alertamiento automático, la exportación de reportes y la centralización de la información, se validaron los siete factores de riesgo y se identificó la clasificación binaria como limitación frente al objetivo de predicción continua, con implicaciones directas sobre la arquitectura. En la Fase 2 se validó el conjunto de datos sintético, se verificó el cumplimiento del objetivo analítico mediante el MLP y se confirmó que la transición hacia la organización modular del backend, junto con el mecanismo de contingencia y la refactorización del modelo de datos, constituyó la condición técnica necesaria para sostener las nuevas capacidades predictivas y garantizar la continuidad operativa. La verificación del cumplimiento de los objetivos, articulada en la Tabla 6.3 y contrastada con los criterios de éxito C1–C5, permite afirmar que el objetivo general fue alcanzado en el contexto del caso de estudio. Este resultado se sitúa, finalmente, dentro de las limitaciones declaradas: una validación cualitativa, en un caso único y apoyada en datos sintéticos para la Fase 2, cuya consolidación empírica constituye la línea prioritaria de trabajo futuro.

Conclusiones

7.1. Conclusiones

El trabajo partió de una necesidad concreta del caso de estudio: la gestión fragmentada y no estandarizada de la información de rendimiento en una academia de natación formativa, que dificultaba el seguimiento, la detección temprana del bajo rendimiento y la comunicación entre entrenador, deportistas y acudientes. Frente a ese problema, el proyecto demostró que es viable materializar un prototipo de plataforma web modular que centraliza los datos, integra un modelo analítico, genera alertas tempranas y apoya la toma de decisiones, según se verificó en Fabio Toro Academy. El aporte central, por tanto, no reside en una mejora deportiva medida, expresamente fuera del alcance, sino en una solución de ingeniería de software que articula gestión de datos, aprendizaje automático, arquitectura modular y validación en un entorno real. Las conclusiones que siguen interpretan ese logro a la luz de cada objetivo específico.

Sobre la identificación y caracterización de los datos

El primer objetivo se cumplió no solo al caracterizar las variables de carga externa, carga interna y bienestar subjetivo y construir un *pipeline* de preprocesamiento trazable (Anexo A), sino al evidenciar que, en un contexto de formación deportiva, la calidad y la adherencia al registro constituyen una condición previa a cualquier modelado. La baja adherencia al formulario, consistente con lo reportado sobre adopción tecnológica en academias de contextos similares (Garcia et al., 2023), no se ocultó ni se sustituyó por supuestos: derivó en la construcción de un conjunto de datos sintético contextualizado, anclado en la programación real del semestre y en el historial deportivo de los atletas. La conclusión relevante es que el proyecto transformó una limitación de datos en una decisión metodológica explícita, cuyo valor es de viabilidad técnica y no de evidencia empírica definitiva.

Sobre el modelo analítico y predictivo

El segundo objetivo se alcanzó cubriendo un espectro analítico que va de la clasificación de bajo rendimiento (Fase 1) a la predicción continua del tiempo esperado en la sesión siguiente mediante un Perceptrón Multicapa (Fase 2), con las métricas reportadas en el Capítulo 4 y sustentadas en el Anexo B. Más que las cifras, la conclusión metodológicamente significativa es doble. Primero, que la propia validación reveló que una salida binaria no satisfacía en sentido estricto el objetivo de

estimación individual del desempeño, lo que motivó la reorientación hacia el modelo de regresión. Segundo, que la concentración de la importancia en variables de bienestar y dolor no puede leerse aún como un hallazgo sobre los determinantes del rendimiento, porque varias de esas variables intervienen en la definición de la etiqueta y producen una fuga de información reconocida; constituye, por tanto, una hipótesis a verificar con datos reales y sin las variables constitutivas del objetivo.

Sobre el diseño y desarrollo del prototipo

El tercer objetivo se materializó en un prototipo funcional que sustituye la gestión dispersa por una plataforma centralizada, consultable e integrada con el componente analítico, con alertas y reportes orientados a los distintos actores. La conclusión de ingeniería más sólida es que la arquitectura evolucionó de un monolito por capas a un Monolito Modular basado en *Domain-Driven Design* de manera *gradual y justificada*, en respuesta a requerimientos y atributos de calidad concretos y no a tendencias tecnológicas. La trazabilidad explícita entre requerimientos y decisiones arquitectónicas, junto con la incorporación de un mecanismo de contingencia que preserva el registro de la sesión ante fallos del servicio analítico, demuestran que la modularidad y la sostenibilidad técnica se abordaron como objetivos de diseño y no como atributos accidentales (Guntakandla, 2025).

Sobre la validación en el caso de estudio

El cuarto objetivo se cumplió mediante una validación con diseño explícito, articulada en dimensiones funcional, técnica, de insumo de datos y de utilidad percibida, ejecutada de forma iterativa. Su conclusión más valiosa es que la validación operó como mecanismo formativo y no como trámite final: la evaluación técnica de la Fase 1 reorientó el desarrollo, y la Fase 2 verificó que las mejoras respondieran a las brechas detectadas. La comparación cualitativa del proceso antes y después confirmó la utilidad percibida del sistema para el seguimiento, el alertamiento y la comunicación entre actores (Kaya and Senbel, 2023), dentro de los límites propios de un caso de estudio único.

Limitaciones de la investigación

Las conclusiones anteriores deben leerse dentro de límites precisos. La validación se realizó en un caso de estudio único y se apoyó en el juicio de un entrenador jefe y de un experto, sin instrumentos cuantitativos de usabilidad ni verificación formal por requerimiento. Las métricas predictivas de la Fase 2 se obtuvieron sobre datos sintéticos y con la alta correlación entre el tiempo planificado y la variable objetivo, factores que pueden inflar el desempeño aparente y ocultar el aporte de las variables de bienestar y carga. El tamaño limitado del dataset obtenido en la Fase 1, la baja adherencia al registro y la ausencia de evaluación longitudinal restringen la generalización de los hallazgos. Estas limitaciones no contradicen el cumplimiento de los objetivos, pero delimitan su alcance y definen la agenda de trabajo futuro.

7.2. Trabajos futuros

Con base en los resultados obtenidos y las limitaciones identificadas, se proponen las siguientes líneas de trabajo para etapas posteriores del proyecto.

Datos y modelo analítico

- **Recalibración con datos reales y *pipeline* de reentrenamiento.** La prioridad es sustituir progresivamente el conjunto sintético por datos empíricos de la academia y formalizar un proceso de reentrenamiento periódico que excluya las variables constitutivas del objetivo, corrija el sesgo de sobreestimación identificado y permita verificar si el peso predictivo del bienestar y el dolor se sostiene en condiciones reales.
- **Estrategias de adherencia a la captura.** Dado que la continuidad del modelo depende del registro sostenido, se recomienda explorar mecanismos integrados al flujo de entrenamiento, como notificaciones automáticas e interfaces simplificadas para dispositivos móviles, que reduzcan la fricción de diligenciamiento (Garcia et al., 2023).
- **Incorporación de variables objetivas.** La integración futura de métricas objetivas provenientes de dispositivos portátiles de bajo costo podría enriquecer el poder predictivo del modelo y reducir su dependencia de la autopercepción del atleta.
- **Interpretabilidad del modelo.** La aplicación de técnicas de inteligencia artificial explicable, como SHAP o LIME, sobre el modelo de regresión permitiría que las predicciones continuas sean tan accionables e interpretables como las alertas ya validadas por el cuerpo técnico.

Arquitectura, despliegue y operación

- **Despliegue en la nube y plan de costos para producción.** El paso del prototipo a un entorno productivo requiere un plan de despliegue en infraestructura de nube acompañado de un análisis de costos operativos, condición necesaria para una adopción sostenible más allá del caso de estudio.
- **Escalamiento selectivo hacia microservicios.** La organización en Monolito Modular basada en DDD habilita una ruta natural para externalizar como servicios independientes los módulos de mayor demanda.
- **Monitoreo, observabilidad y trazabilidad.** Sobre la base del estado de procesamiento de sesiones ya incorporado, se sugiere instrumentar registros, métricas de salud del servicio analítico y trazabilidad del ciclo de vida de cada registro, para sostener la mantenibilidad y diagnosticar incidencias en operación.

Producto y validación

- **Verificación formal de requerimientos y usabilidad.** Cerrar la limitación señalada en la validación mediante un protocolo de verificación requerimiento por requerimiento e instrumentos cuantitativos de usabilidad con los distintos perfiles de usuario.
- **Replicación en otras academias.** Extender la implementación a otras academias de natación formativa permitiría evaluar la validez externa de la solución e identificar ajustes dependientes del contexto.
- **Módulo de comunicación directa entre actores.** Un componente de mensajería o notificaciones internas cerraría de forma más ágil y trazable el ciclo comunicativo que hoy se apoya en reportes y tableros ([College of Education and Human Development, 2023](#)).

7.3. Lecciones aprendidas

El desarrollo del proyecto dejó aprendizajes en los planos técnico, metodológico y humano que trascienden los resultados puntuales.

- **La calidad de los datos antecede a cualquier modelo.** El hallazgo más temprano y determinante fue que los datos disponibles eran insuficientes para entrenar directamente los modelos. Esta restricción evidenció que la madurez analítica de una organización deportiva depende, antes que de los algoritmos, de sus prácticas de registro y estandarización.
- **La validación temprana con usuarios reorienta decisiones técnicas.** La evaluación del componente analítico durante la Fase 1 fue el detonante de la migración hacia el modelo de regresión. La retroalimentación no debe posponerse al cierre: anticiparla ahorra rediseños costosos y mejora la pertinencia del sistema.
- **La arquitectura debe crecer con el sistema.** La evolución hacia el Monolito Modular se justificó por necesidades concretas surgidas en la validación; adoptar microservicios desde el inicio habría introducido una complejidad desproporcionada. El sobrediseño arquitectónico temprano puede ser tan perjudicial como la ausencia de diseño.
- **La adopción tecnológica es un proceso social, no solo técnico.** La baja adherencia al registro mostró que disponer de una herramienta funcional no garantiza su uso sostenido: la incorporación de tecnología en entornos formativos requiere acompañamiento, capacitación y comunicación de beneficios.
- **El alcance realista y la honestidad metodológica sostienen la confianza.** Delimitar explícitamente lo que el prototipo puede y no puede hacer, y comunicar con transparencia las limitaciones del conjunto sintético, fue decisivo para mantener expectativas alineadas con los actores del caso de estudio.

- **La interdisciplinariedad es un activo.** Integrar ciencias del deporte, ingeniería de software y aprendizaje automático exigió un esfuerzo de traducción conceptual permanente, pero produjo un sistema más pertinente, cuyas decisiones técnicas están ancladas en el conocimiento del dominio y no solo en criterios estadísticos.

Bibliografía

- Analytics Vidhya (2024). Distinguish between tree-based machine learning models. <https://www.analyticsvidhya.com/blog/2021/04/distinguish-between-tree-based-machine-learning-algorithms/>.
- AthleteMonitoring.com (2024). Athlete management system for youth sports.
- Aydemir, B., Aydoğan, M. T., Boz, E., Kul, M., Kirkbir, F., and Özkara, A. B. (2025). Validity and reliability of a novel ai-based system in athletic performance assessment: The case of deepsport. *Sensors*, 25:5580.
- Aydın, A., Özgünen, K. T., Keskinoglu, C., Talib, E. N., Dağ, N. Y., and Gezgin, E. (2025). Development of a low-cost team performance tracking system. *BMC Sports Science, Medicine and Rehabilitation*, 17:281.
- Barry, L., Lyons, M., McCreesh, K., Myers, T., Powell, C., and Comyns, T. (2024). The relationship between training load and injury in competitive swimming: A two-year longitudinal study. *Applied Sciences (Switzerland)*, 14.
- Benjamin, M. (2025). Microservices architecture for scalable enterprise applications. Technical report, ResearchGate.
- Bernaciková, M., Kumstát, M., Burešová, I., Kapounková, K., Struhár, I., Sebera, M., and Paludo, A. C. (2022). Preventing chronic fatigue in czech young athletes: The features description of the “smartraining” mobile application. *Frontiers in Physiology*, 13.
- Bianchi, V., Ambrosini, L., Presta, V., Gobbi, G., and Munari, I. D. (2022). Prediction of kick count in triathletes during freestyle swimming session using inertial sensor technology. *Applied Sciences (Switzerland)*, 12.
- Bischl, B., Binder, M., Lang, M., Pielok, T., Richter, J., Coors, S., Thomas, J., Ullmann, T., Becker, M., Boulesteix, A.-L., Deng, D., and Lindauer, M. (2023). Hyperparameter optimization: Foundations, algorithms, best practices, and open challenges. *WIREs Data Mining and Knowledge Discovery*, 13.
- Bittencourt, N. F. N., Meeuwisse, W. H., Mendonça, L. D., Nettel-Aguirre, A., Ocarino, J. M., and Fonseca, S. T. (2016). Complex systems approach for sports injuries: moving from risk factor identification to injury pattern recognition—narrative review and new concept. *British Journal of Sports Medicine*, 50:1309–1314.
- Carvalho, D. D., Goethel, M. F., Silva, A. J., Vilas-Boas, J. P., Pyne, D. B., and Fernandes, R. J. (2024). Swimming performance interpreted through explainable artificial intelligence (xai)—practical tests and training variables modelling. *Applied Sciences (Switzerland)*, 14.

- Chamari, L., Walker, S., Petrova, E., and Pauwels, P. (2025). Towards portable model predictive control-based applications for demand side management in buildings. *Energy and Buildings*, 347.
- Cheng, L. (2022). Implementation of snow and ice sports health and sports information collection system based on internet of things. *Journal of Healthcare Engineering*, 2022:1–12.
- Chicco, D. and Jurman, G. (2020). The advantages of the matthews correlation coefficient (mcc) over f1 score and accuracy in binary classification evaluation. *BMC Genomics*, 21:6.
- College of Education and Human Development (2023). How does the relationship between parents and coaches affect young athletes?
- Contreras, D. A. B. (2020). Arquitectura de microservicios. *Tecnología, Investigación y Academia TIA*.
- Cordeiro, M. C., Cathain, C. O., Daly, L., Kelly, D. T., and Rodrigues, T. B. (2025). A synthetic data-driven machine learning approach for athlete performance attenuation prediction. *Frontiers in Sports and Active Living*, 7.
- de la Hoz, J. (2025). Colombia consolida su reserva deportiva con una política estructural y resultados históricos. <https://www.mindeporte.gov.co/sala-prensa/noticias-mindeporte/colombia-consolida-reserva-deportiva-politica-estructural-resultados-historicos>.
- Eetvelde, H. V., Mendonça, L. D., Ley, C., Seil, R., and Tischer, T. (2021). Machine learning methods in sport injury prediction and prevention: a systematic review. *Journal of Experimental Orthopaedics*, 8.
- FECNA (2024). Instructivo oficial federación colombiana de natación 2024. Technical report, Federación Colombiana de Natación.
- Foster, C., Rodriguez-Marroyo, J. A., and de Koning, J. J. (2017). Monitoring training loads: The past, the present, and the future. *International Journal of Sports Physiology and Performance*, 12:S2–2–S2–8.
- Gabbett, T. J. (2016). The training—injury prevention paradox: should athletes be training smarter and harder? *British Journal of Sports Medicine*, 50:273–280.
- Garcia, C., Suarez, J., Niño, D., Roa, O., Barbosa, A., Plazas, J., Garcia, J., López, J., and Picón, I. (2023). Data as a tool for regulation and decision-making in telecommunications services: An academy-state innovation experience. In *Proceedings of the International Conference on Data and Telecommunications Regulation*, pages 1–14.
- Gonzalez, R. A., Giménez, S. A., Molina, P., Numa, R., and Zalazar, R. (2024). ¿son los microservicios la mejor opción? una evaluación de su eficacia y eficiencia frente a los monolitos. Technical report, Universidad Nacional de Misiones.

- Goudsmit, J., Otter, R. T., Stoter, I., van Holland, B., van der Zwaard, S., de Jong, J., and Vos, S. (2022). Co-operative design of a coach dashboard for training monitoring and feedback. *Sensors*, 22.
- Guntakandla, A. R. (2025). Modular architecture: A scalable and efficient system design approach for enterprise applications. *World Journal of Advanced Research and Reviews*, 26:3114–3126.
- He, G. and Choi, H. S. (2025). Stacked ensemble model for NBA game outcome prediction analysis. *Scientific Reports*, 15:29983.
- Jiang, T., Gradus, J. L., and Rosellini, A. J. (2020). Supervised machine learning: A brief primer. *Behavior Therapy*, 51:675–687.
- Kaya, T. and Senbel, S. (2023). A dynamic online dashboard for tracking the performance of division 1 basketball athletic performance. Technical report, Sacred Heart University.
- Kellmann, M., Bertollo, M., Bosquet, L., Brink, M., Coutts, A. J., Duffield, R., Erlacher, D., Halson, S. L., Hecksteden, A., Heidari, J., Kallus, K. W., Meeusen, R., Mujika, I., Robazza, C., Skorski, S., Venter, R., and Beckmann, J. (2018). Recovery and performance in sport: Consensus statement. *International Journal of Sports Physiology and Performance*, 13:240–245.
- Luo, Y., Quan, T., and Cao, Y. (2025). Predicting football match outcomes: a multilayer perceptron neural network model based on technical statistics indicators of the fifa world cup. *Frontiers in Sports and Active Living*, 7:1705198.
- López, D. and Maya, E. (2024). Arquitectura de software basada en microservicios para desarrollo de aplicaciones web. Technical report, Asamblea Nacional del Ecuador.
- López-Valenciano, A., Ayala, F., Puerta, J. M., Croix, M. B. A. D. S., Vera-Garcia, F. J., Hernández-Sánchez, S., Ruiz-Pérez, I., and Myer, G. D. (2018). A preventive model for muscle injuries. *Medicine & Science in Sports & Exercise*, 50:915–927.
- Mindeporte (2025). En 2025, colombia consolida su trabajo en la reserva deportiva con inversión y formación de nuevos talentos. <https://www.mindeporte.gov.co/sala-prensa/noticias-mindeporte/2025-colombia-consolida-trabajo-reserva-deportiva-inversion-formacion-nuevos-talentos>.
- MinTIC (2025). Documento maestro de los lineamientos del modelo de seguridad y privacidad de la información. Technical report, Ministerio de Tecnologías de la Información y las Comunicaciones.
- MultimediaEval (2022). Swimtrack: Swimmers and stroke rate detection in elite race videos.
- Penggalih, M. H. S. T., Trisnantoro, L., Prabandari, Y. S., Surono, S., Rahadian, B., Ariyanto, M. A., Niamilah, I., Solichah, K. M., Sahdani, F., Rizana, N. A., Danumaya, A., Muslichah, R., Dewinta, M. C. N., Reswati, V. D. Y., and Bactiar, N. (2025). Exploring sport science

- implementation in indonesian youth athlete training centers: A consolidated framework of implementation research-based thematic analysis. *International Journal of Sports Science & Coaching*, 20:1876–1888.
- Qian, H. and Lee, S. (2025). A multidimensional prediction model for overtraining risk in youth soccer players: Integrating physiological and psychological markers. *Journal of Sports Sciences*, 43:1819–1834.
- Saw, A. E., Main, L. C., and Gustin, P. B. (2016). Monitoring the athlete training response: subjective self-reported measures trump commonly used objective measures: a systematic review. *British Journal of Sports Medicine*, 50:281–291.
- Shetty, B. (2023). 5 classification algorithms for machine learning. <https://builtin.com/data-science/supervised-machine-learning-classification>.
- SwimCloud (2025). Meet management & team platform.
- van Buuren, S. (2018). *Flexible Imputation of Missing Data, Second Edition*. Chapman and Hall/CRC.
- Yang, L. and Shami, A. (2020). On hyperparameter optimization of machine learning algorithms: Theory and practice. *Neurocomputing*, 415:295–316.
- Zhang, W., Shen, X., Zhang, H., et al. (2024). Feature importance measure of a multilayer perceptron based on the presingle-connection layer. *Knowledge and Information Systems*, 66:511–533.

Detalle de la recolección, el preprocesamiento y el conjunto de datos

Este anexo amplía la recolección y el preprocesamiento de los datos reales (Fase 1) y detalla la construcción y el análisis del conjunto sintético (Fase 2).

A.1. Fase 1: detalle de la recolección y el preprocesamiento

A.1.1. Justificación del proceso de recolección

El formulario se concibió bajo el principio de *captura oportuna*: completar el registro dentro de los treinta minutos posteriores a la sesión reduce el sesgo de recuerdo que afecta a las medidas subjetivas, especialmente al RPE (Saw et al., 2016). La estandarización del mismo garantiza que todos los nadadores reporten las mismas variables con las mismas escalas, mientras que la centralización en una única hoja de respuestas y la accesibilidad desde un dispositivo móvil, sin hardware especializado, favorecen la continuidad del registro en el contexto de una academia de formación (Foster et al., 2017; Garcia et al., 2023). Estas decisiones, no obstante, no eliminan la dependencia de la disciplina de reporte de los atletas, que, como evidencia la caracterización del corpus, resultó ser el principal factor limitante de la Fase 1.

A.1.2. Diccionario completo de variables

La Tabla A.1 describe todas las variables capturadas y derivadas, con su tipo, escala y rol previsto en el modelado.

Tabla A.1: Diccionario de variables del conjunto de datos.

Variable	Descripción	Tipo	Escala / unidad
fecha	Fecha de la sesión	Fecha	dd/mm/aaaa
nadador	Identificador del atleta	Categórica	Anonimizado
asistió	Asistencia a la sesión	Binaria	Sí / No
volumen_metros	Volumen total nadado	Numérica	Metros
zona_aerobica_min	Minutos en zona aeróbica	Numérica	Minutos
zona_anaerobica_min	Minutos en zona anaeróbica	Numérica	Minutos
tipo_sesion	Tipo de sesión	Categórica	R1, R2, R1-2, VO2, RL, TL
tiempo_esperado	Tiempo de 100 m planificado (equivalente a <code>tiempo_100m_plan_sec</code> en la Fase 2)	Numérica	Segundos
tiempo_real	Tiempo de 100 m logrado (equivalente a <code>tiempo_100m_real_sec</code> en la Fase 2)	Numérica	Segundos
rpe	Percepción del esfuerzo	Ordinal	1 a 10
horas_sueno	Horas de sueño previas	Numérica	Horas
estado_animo	Estado de ánimo	Ordinal	Muy mal / Normal / Excelente
nivel_estres	Nivel de estrés	Ordinal	1 a 5
motivacion	Motivación para entrenar	Ordinal	1 a 10
dolor_presente	Presencia de dolor o molestia	Binaria	0 / 1
dias_desde_competencia	Días desde la última competencia	Numérica	Días
fase_temporada	Fase del macrociclo	Categórica	Pretemporada, Preparación, Taper, Postcompetencia
<i>Variables derivadas (Etapa 7)</i>			
srpe	Carga interna de la sesión	Numérica	UA (Ecuación 3.1)
ratio_aerobico	Polarización del entrenamiento	Numérica	[0, 1] (Ecuación 3.2)
indice_bienestar	Bienestar integrado (normalizado)	Numérica	[0, 1] (Ec. 3.3)

A.1.3. Detalle de las etapas de preprocesamiento

A continuación se describe cada etapa del pipeline, su propósito y su comportamiento observado sobre los datos reales.

1. **Estandarización de nomenclatura.** Se renombran las columnas del formulario a `snake_case`

mediante un mapeo explícito, lo que evita ambigüedades en el código posterior.

2. **Normalización de formatos numéricos.** El volumen nadado puede llegar con separadores heterogéneos (por ejemplo, “4.500”, “4,500” o “4500”); se aplican reglas de dominio para unificarlo y convertirlo a un valor numérico.
3. **Conversión de tiempos.** Los tiempos de 100 m se llevan a segundos a partir de los distintos formatos de entrada (por ejemplo, “1:25” $\rightarrow$ 85 s).
4. **Codificación de variables categóricas.** El estado de ánimo se codifica de forma ordinal (muy mal = 1, normal = 2, excelente = 3); el dolor y la asistencia, de forma binaria; y la fase del macrociclo se mapea a sus categorías. Las respuestas no reconocidas en este mapeo se etiquetan como *Desconocido*, lo que explica el 4,2 % de sesiones sin fase clasificada.
5. **Validación de rangos y fechas.** Se definen *y se aplican* los límites fisiológicamente plausibles de cada variable (RPE, sueño, volumen y tiempos de 100 m): los valores fuera de rango se marcan como nulos. De este modo, los errores de digitación (por ejemplo, tiempos de 100 m fuera del intervalo plausible o fechas imposibles) ya no sobreviven al preprocesamiento.
6. **Tratamiento de valores ausentes.** Se distinguen los nulos legítimos (por ejemplo, una sesión sin asistencia no tiene datos de entrenamiento) de la información genuinamente faltante. Las sesiones sin asistencia (25 de 216) se separan del conjunto de entrenamiento.
7. **Creación de variables derivadas.** Se calculan el sRPE, el ratio aeróbico y el índice de bienestar según las Ecuaciones 3.1 a 3.3.

A.1.4. Resultados del preprocesamiento sobre los datos reales

A.1.4.1. Estadísticos descriptivos

La Tabla A.2 presenta los estadísticos de las variables numéricas sobre las 191 sesiones con asistencia.

Tabla A.2: Estadísticos descriptivos de las variables principales (datos reales, $n = 191$).

Variable	Media	DE	Mín	Máx
Volumen (m)	4896,3	1381,2	1100,0	10000,0
Zona aeróbica (min)	33,6	16,6	1,0	75,0
Zona anaeróbica (min)	38,7	15,3	0,0	80,0
RPE	5,6	1,6	1,0	10,0
Horas de sueño	7,5	0,9	5,0	10,0
Nivel de estrés	1,9	1,0	1,0	5,0
Motivación	6,3	2,1	1,0	10,0
sRPE (UA)	417,8	187,2	11,6	960,0
Ratio aeróbico	0,46	0,18	0,01	1,00
Índice de bienestar	0,703	0,273	0,000	1,000

A.1.4.2. Distribución por nadador

La Tabla A.3 confirma la fuerte concentración de los registros: el conjunto está dominado por un atleta (37,2 %) y cuatro nadadores aportan una sola sesión.

Tabla A.3: Sesiones por nadador (etiquetas anonimizadas, datos reales).

Nadador	Sesiones	%
Nadador 1	71	37,2
Nadador 2	31	16,2
Nadador 3	26	13,6
Nadador 4	13	6,8
Nadador 5	12	6,3
Nadador 6	12	6,3
Nadador 7	10	5,2
Nadador 8	7	3,7
Nadador 9	5	2,6
Nadador 10	1	0,5
Nadador 11	1	0,5
Nadador 12	1	0,5
Nadador 13	1	0,5
Total	191	100

A.1.4.3. Distribución de variables categóricas

La Tabla A.4 resume las distribuciones categóricas. Destacan el predominio de sesiones de baja intensidad (R1 y R1-2 suman cerca del 77 %), la ausencia de sesiones TL y la concentración del estado de ánimo en *normal*.

Tabla A.4: Distribución de variables categóricas (datos reales, $n = 191$).

Categoría	Frecuencia (%)
<i>Tipo de sesión</i>	
R1	80 (41,9)
R1-2	67 (35,1)
R2	20 (10,5)
RL	14 (7,3)
VO2	10 (5,2)
<i>Fase del macrociclo</i>	
Preparación	95 (49,7)
Pretemporada	52 (27,2)
Postcompetencia	36 (18,8)
Desconocido	8 (4,2)
<i>Estado de ánimo</i>	
Muy mal	12 (6,3)
Normal	105 (55,0)
Excelente	74 (38,7)
<i>Presencia de dolor</i>	
Sin dolor	141 (73,8)
Con dolor	50 (26,2)

A.2. Fase 2: detalle de la construcción y el análisis del conjunto sintético

A.2.1. Supuestos de generación

La generación se sustenta en supuestos mínimos, declarados de forma explícita y validados con el entrenador jefe. Su formulación, contrastada con la implementación del cuaderno, es la siguiente:

1. **Reproducibilidad.** La semilla aleatoria se fija en 42, de modo que toda ejecución del cuaderno produce el mismo conjunto.
2. **Asistencia.** La asistencia diaria es alta (cercana al 92 %); las sesiones sin asistencia no generan datos de entrenamiento.

3. **Cronología deportiva.** Las sesiones se anclan al macrociclo de diecisiete semanas de la academia (Anexo C), con progresión de capacidades e incremento del volumen a lo largo de la temporada.
4. **Intensidad por tipo de sesión.** El tiempo planificado se modula con factores de intensidad propios de cada tipo de sesión.
5. **Bienestar sesgado a estados aceptables.** Las variables de bienestar (ánimo, estrés, motivación) se muestrean de distribuciones centradas en valores normales.
6. **Coherencia interna.** Se inducen relaciones acordes con la lógica deportiva: el dolor incrementa el estrés y reduce la motivación, y un mejor bienestar se asocia a mejores tiempos.
7. **Variabilidad controlada.** El tiempo real se obtiene sumando al planificado desviaciones dependientes del estado y un componente aleatorio de baja magnitud, con aproximadamente un 5 % de casos que invierten la relación esperada.

Precisión sobre el dolor. A diferencia de una persistencia del 40 % entre días que pudiera suponerse, en la implementación el dolor se genera de forma independiente por sesión, con probabilidad cercana al 15 %. Esta nota evita atribuir al conjunto una dinámica temporal del dolor que el código no produce.

A.2.2. Índice de bienestar: definición unificada

El índice de bienestar es *único* para todo el trabajo y está normalizado en $[0, 1]$, ponderando el dolor con un 40 %:

$$IB = 0,20 m_{\text{norm}} + 0,20 (1 - e_{\text{norm}}) + 0,20 a_{\text{norm}} + 0,40 d_{\text{num}} \quad (\text{A.1})$$

donde $m_{\text{norm}} = (\text{motivación} - 1)/9$, $e_{\text{norm}} = (\text{estrés} - 1)/4$, $a_{\text{norm}} = (\text{ánimo} - 1)/2 \in \{0; 0,5; 1\}$ y $d_{\text{num}} = 1$ si no hay dolor y 0 en caso contrario.

Esta es exactamente la misma fórmula que la Ecuación 3.3: el preprocesamiento de los datos reales y el generador del conjunto sintético la aplican de forma idéntica. Por ello, el umbral “índice de bienestar $\leq 0,3$ ”, empleado para construir la variable objetivo, representa el mismo nivel de bienestar en ambos conjuntos. La unificación corrige una inconsistencia de versiones anteriores, en las que el preprocesamiento empleaba una definición sin normalizar de rango aproximado entre $-1,1$ y $1,4$, lo que volvía no comparables los umbrales entre fases.

A.2.3. Flujo de construcción y validaciones automáticas

El cuaderno construye el conjunto iterando, por nadador y por día del calendario, los siguientes pasos: (i) se decide la asistencia; (ii) se genera el contexto de bienestar (sueño, estrés, motivación, ánimo, dolor) con las relaciones internas descritas; (iii) se determina el tipo de sesión y el estilo

según la práctica del nadador; (iv) se calcula el volumen y el tiempo de 100 m planificado a partir del tiempo base y de los factores de intensidad; (v) se calcula el tiempo real sumando al planificado las desviaciones por estado, la penalización por dolor y la penalización por sueño insuficiente; y (vi) se derivan el sRPE y el índice de bienestar.

El tiempo real se obtiene como

$$t_{\text{real}} = t_{\text{plan}} + \Delta(\text{estado, dolor, sueño}), \quad (\text{A.2})$$

donde la desviación Δ es positiva y grande cuando el estado es *muy mal* o el índice de bienestar es bajo (por ejemplo, de +6 a +12 s en sesiones R1 o R1-2), puede ser negativa cuando el estado es excelente y no hay dolor, e incorpora penalizaciones aditivas por dolor (de +1 a +3 s) y por dormir cuatro horas o menos. Esta regla es la fuente directa de la asociación inversa fuerte entre bienestar y diferencia de tiempo observada en el análisis. Tras la generación, el cuaderno ejecuta una batería de validaciones automáticas (rangos, ausencia de duplicados indebidos, coherencia de las relaciones inducidas) que verifican la integridad del conjunto antes de exportarlo.

A.2.4. Análisis exploratorio detallado

Las tablas siguientes resumen las distribuciones del conjunto sintético sobre las 1.559 sesiones con asistencia, reproducibles con semilla 42. Las figuras correspondientes se generan con el script entregado.

Tabla A.5: Distribución por estilo y por tipo de sesión (sintético, asistentes, $n = 1559$).

Estilo	n	Tipo de sesión	n
Libre	769	R2	623
Mariposa	244	VO2	313
Pecho	222	R1	237
Espalda	194	RL	157
Combinado	130	TL	156
		R1-2	73

Tabla A.6: Distribución de bienestar y dolor (sintético, asistentes, $n = 1559$).

Categoría	Frecuencia
<i>Estado de ánimo</i>	
Normal	1064
Muy mal	350
Excelente	145
<i>Presencia de dolor</i>	
No	1320
Sí	239

El índice de bienestar sintético toma valores en el intervalo $[0,09; 1,00]$ y la diferencia de tiempo, entre $-5,0$ y $+13,4$ s, con media $+1,41$ s. Estas distribuciones, junto con la correlación $r = -0,70$ entre bienestar y diferencia de tiempo, confirman que el generador reproduce la estructura diseñada. Se reitera que estos resultados verifican el comportamiento del modelo de reglas y no constituyen evidencia empírica del rendimiento de los nadadores.

Figuras del anexo. Las Figuras A.1 a A.6 ilustran, respectivamente: la distribución por estilo (A.1) y por tipo de sesión (A.2); el histograma de la diferencia de tiempo (A.3); la dispersión entre el índice de bienestar y la diferencia de tiempo (A.4); la serie temporal de la diferencia media diaria (A.5); y las distribuciones de estado de ánimo y dolor (A.6).

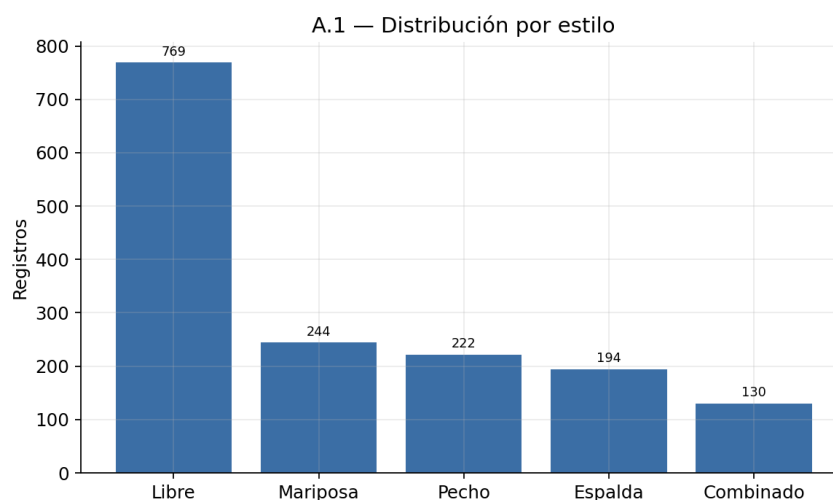


Figura A.1: Distribución de registros por estilo (conjunto sintético).

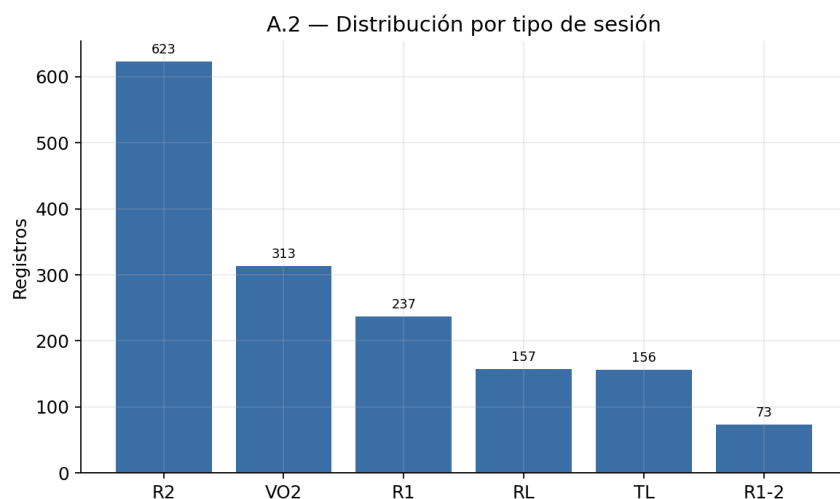


Figura A.2: Distribución de registros por tipo de sesión (conjunto sintético).

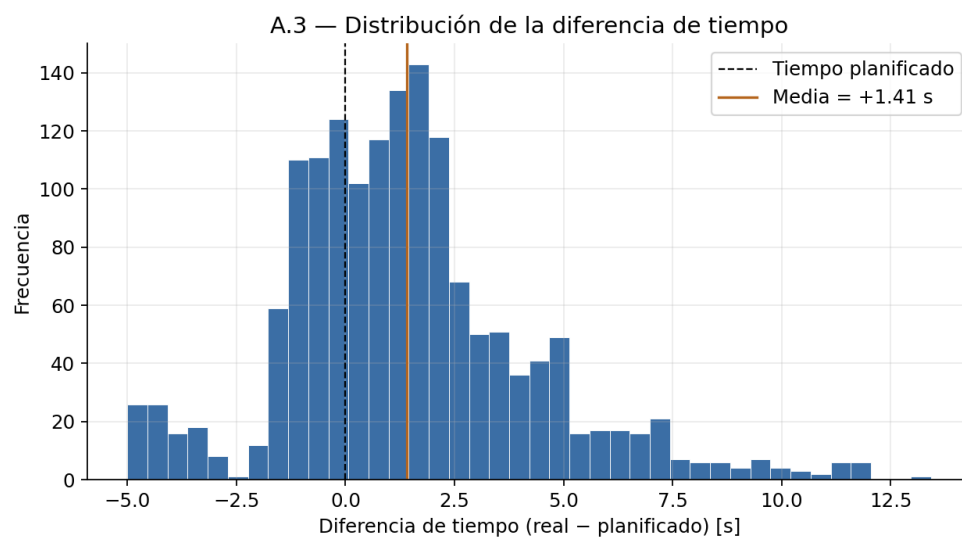


Figura A.3: Distribución de la diferencia de tiempo (tiempo real menos planificado).

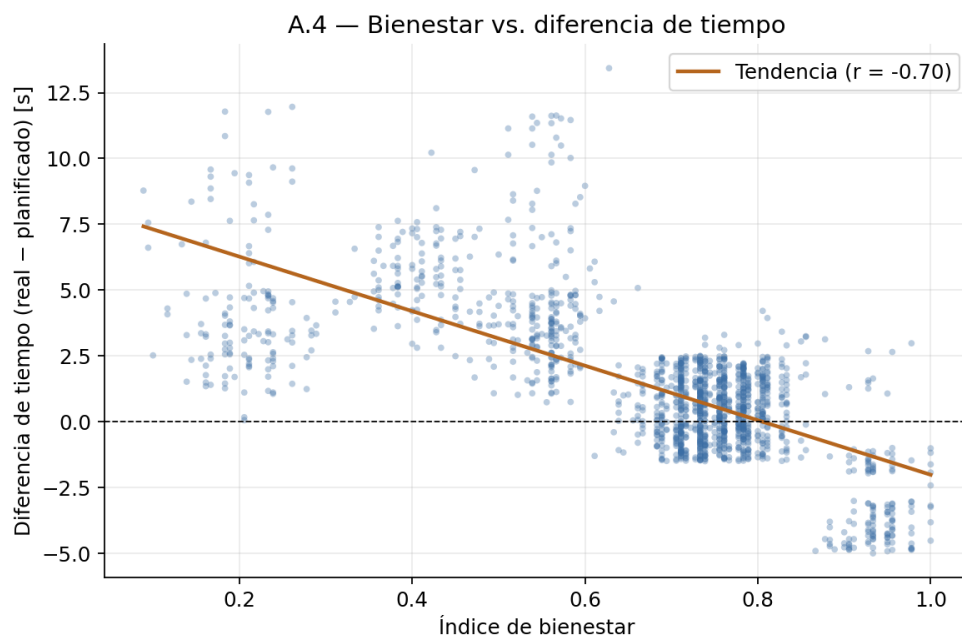


Figura A.4: Relación entre el índice de bienestar y la diferencia de tiempo.

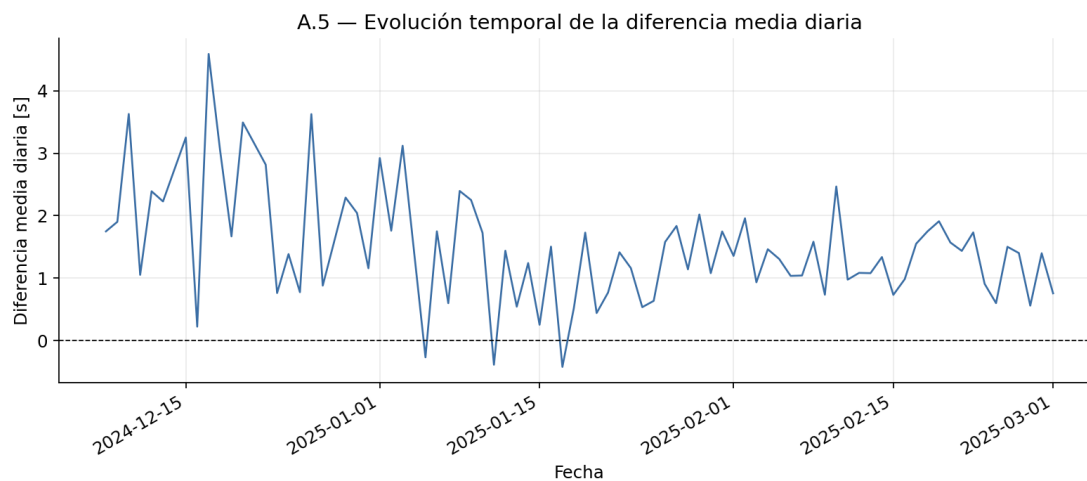


Figura A.5: Evolución temporal de la diferencia media diaria.

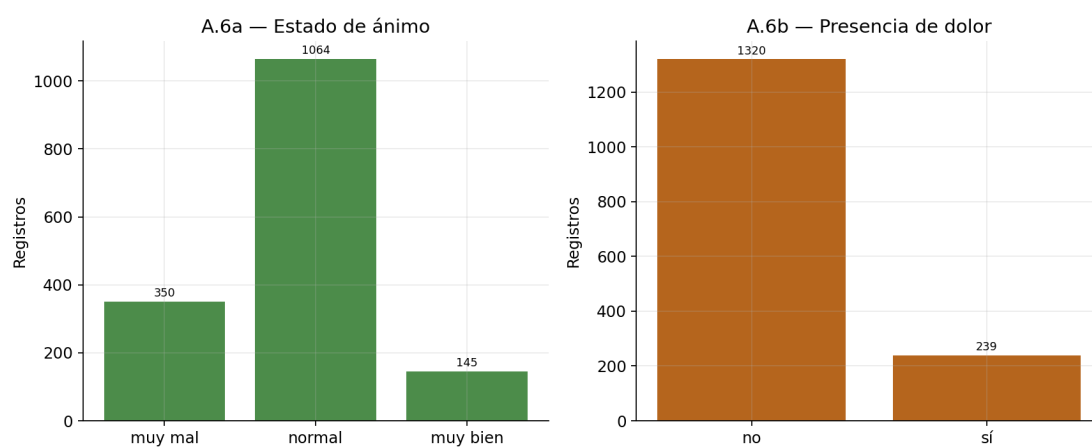


Figura A.6: Distribución de estado de ánimo y presencia de dolor (conjunto sintético).

Detalle técnico del modelado analítico y predictivo

Este anexo reúne el material técnico ampliado del Capítulo 4. Su propósito es documentar los espacios de búsqueda de hiperparámetros, las definiciones formales de las métricas, las visualizaciones completas de la evaluación y las tablas de configuración que sustentan las decisiones tomadas en el cuerpo del capítulo, de modo que cada elección quede explicada.

B.1. Fase 1: detalle del modelo de clasificación de bajo rendimiento

B.1.1. Espacios de búsqueda de hiperparámetros

Los hiperparámetros de cada modelo se optimizaron mediante **GridSearchCV** (búsqueda exhaustiva en rejilla) con **validación cruzada estratificada de 5 folds**, utilizando el F1-Score como métrica de optimización. La validación cruzada divide el conjunto de entrenamiento en 5 subconjuntos, entrena el modelo en 4 de ellos y lo evalúa en el restante, repitiendo el proceso 5 veces para obtener una estimación robusta del rendimiento. El uso del F1-Score en lugar de la exactitud responde a la necesidad de balancear la detección de nadadores en riesgo (*recall*) con la precisión de las alertas, evitando que el modelo se incline a predecir la clase mayoritaria [Chicco and Jurman, 2020](#). La búsqueda de hiperparámetros es un paso fundamental para obtener el mejor rendimiento posible de cada algoritmo sin caer en configuraciones que produzcan sobreajuste [Bischi et al., 2023](#).

B.1.1.1. Regresión Logística

Dado que la Regresión Logística es sensible a la escala de las variables, se aplicó una estandarización previa (*StandardScaler*) que transforma cada variable para que tenga media cero y desviación estándar uno. Esta normalización asegura que todas las variables contribuyan en igualdad de condiciones al proceso de optimización, sin que aquellas con rangos más amplios (como el volumen en metros) dominen sobre las de rango reducido (como el índice de bienestar).

Tabla B.1: Espacio de búsqueda de hiperparámetros para Regresión Logística.

Hiperparámetro	Valores explorados
C (inverso de regularización)	0.001, 0.01, 0.1, 1, 10
Tipo de penalización	L1 (Lasso), L2 (Ridge)
Balanceo de clases	<i>balanced</i> , <i>None</i>
Solver	<i>saga</i> (compatible con L1 y L2)

El parámetro C controla la intensidad de la regularización: valores pequeños imponen una regularización más fuerte, simplificando el modelo a costa de posible subajuste, mientras que valores altos permiten mayor complejidad. La opción `class_weight="balanced"` ajusta automáticamente el peso de cada clase en proporción inversa a su frecuencia, compensando el desbalance entre sesiones normales y de bajo rendimiento.

B.1.1.2. Random Forest

Random Forest no requiere estandarización de las variables, ya que los árboles de decisión basan sus divisiones en umbrales relativos dentro de cada variable. El espacio de hiperparámetros explorado fue más amplio, abarcando tanto la estructura del bosque como los controles de profundidad individual de cada árbol.

Tabla B.2: Espacio de búsqueda de hiperparámetros para Random Forest.

Hiperparámetro	Valores explorados
Número de árboles (<i>n_estimators</i>)	100, 200, 300
Profundidad máxima (<i>max_depth</i>)	3, 5, 7, <i>None</i> (sin límite)
Mínimo de muestras para dividir	5, 10, 15
Mínimo de muestras por hoja	3, 5, 8
Variables por división (<i>max_features</i>)	<i>sqrt</i> , <i>log2</i>
Balanceo de clases	<i>balanced</i> , <i>balanced_subsample</i>

Los parámetros `min_samples_split` y `min_samples_leaf` actúan como controles de sobreajuste al impedir que los árboles crezcan demasiado y memoricen los datos de entrenamiento. La restricción `max_features` limita el número de variables que cada árbol considera en cada división, lo que incrementa la diversidad del bosque y mejora la generalización.

B.1.1.3. XGBoost

XGBoost recibió la configuración más exhaustiva de hiperparámetros, dado que ofrece múltiples mecanismos de control del sobreajuste que deben calibrarse cuidadosamente, especialmente con

conjuntos de datos de tamaño moderado. Se incorporó el parámetro `scale_pos_weight`, calculado como la razón entre el número de muestras de la clase mayoritaria y la minoritaria, para compensar el desbalance de clases directamente en la función de pérdida del gradiente.

Tabla B.3: Espacio de búsqueda de hiperparámetros para XGBoost.

Hiperparámetro	Valores explorados
Tasa de aprendizaje (<i>learning_rate</i>)	0.01, 0.05, 0.1
Número de árboles (<i>n_estimators</i>)	100, 200, 300
Profundidad máxima (<i>max_depth</i>)	3, 5, 7
Peso mínimo por hoja (<i>min_child_weight</i>)	3, 5, 10
Submuestreo de filas (<i>subsample</i>)	0.7, 0.8, 1.0
Submuestreo de columnas (<i>colsample_bytree</i>)	0.7, 0.8, 1.0
Regularización L1 (<i>reg_alpha</i>)	0.0, 0.1, 1.0
Regularización L2 (<i>reg_lambda</i>)	1.0, 2.0, 5.0

Adicionalmente, se implementó *early stopping* con una paciencia de 15 rondas durante la búsqueda de hiperparámetros, lo que detiene el entrenamiento de cada configuración cuando la métrica de validación deja de mejorar, evitando iteraciones innecesarias y reduciendo el riesgo de sobreajuste. Para la validación cruzada final del mejor modelo, se congeló el número óptimo de árboles obtenido durante el *early stopping* y se desactivó este mecanismo, garantizando una evaluación limpia y reproducible.

B.1.2. Definición formal de las métricas de evaluación

Para interpretar adecuadamente los resultados de la Fase 1 es necesario definir cada métrica en el contexto del problema, tomando como referencia las 28 sesiones del conjunto de prueba (18 “Normal” y 10 “Bajo Rendimiento”) y el desempeño del modelo seleccionado, XGBoost.

- **Exactitud (*accuracy*)**. Proporción total de predicciones correctas: de las 28 sesiones evaluadas, XGBoost clasificó correctamente 23 (82.1 %). Esta métrica puede ser engañosa con clases desbalanceadas, ya que un modelo que siempre predijera “Normal” obtendría una exactitud del 64.3 % sin detectar a ningún nadador en riesgo.
- **Precisión (*precision*)**. De todos los nadadores marcados como bajo rendimiento, cuántos lo estaban realmente. XGBoost alcanzó 0.778, es decir, cuando la plataforma emite una alerta acierta en aproximadamente 8 de cada 10 casos; las 2 alertas restantes son falsas alarmas.
- **Recall (sensibilidad)**. De todos los nadadores que realmente estaban en bajo rendimiento, cuántos detectó el modelo. XGBoost detectó 7 de los 10 (*recall* = 0.70); los 3 no detectados son falsos negativos. En un contexto deportivo juvenil, donde la salud del atleta es prioritaria,

esta es una métrica particularmente sensible, ya que el costo de un falso negativo (un nadador en riesgo sin intervención) supera al de un falso positivo (una revisión innecesaria).

- **F1-Score.** Media armónica entre precisión y *recall*; solo es alta cuando ambas lo son simultáneamente. XGBoost obtuvo el mejor valor (0.7368), reflejando el mejor compromiso entre minimizar falsas alarmas y maximizar la detección [Chicco and Jurman, 2020](#).
- **AUC-ROC.** Mide la capacidad de discriminar entre clases independientemente del umbral de decisión: 0.5 equivale a un clasificador aleatorio y 1.0 a discriminación perfecta. Regresión Logística y Random Forest comparten un AUC de 0.917, superior al 0.869 de XGBoost, lo que indica que producen probabilidades mejor calibradas, aunque XGBoost genera mejores decisiones finales con el umbral estándar de 0.5.

B.1.3. Matrices de confusión

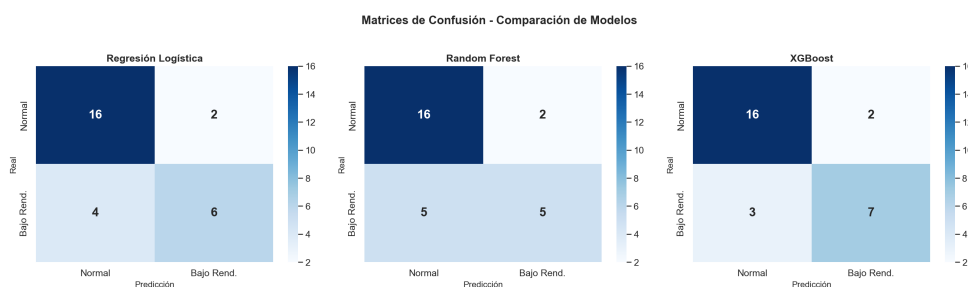


Figura B.1: Matrices de confusión para los tres modelos evaluados en el conjunto de prueba.

Las matrices de confusión visualizan en detalle las predicciones de cada modelo. El eje vertical representa la clase real y el horizontal la predicción; la celda superior izquierda corresponde a los verdaderos negativos y la inferior derecha a los verdaderos positivos. Los tres modelos muestran un comportamiento idéntico en sesiones normales (16 verdaderos negativos y 2 falsos positivos). La diferencia sustancial está en la detección de bajo rendimiento: Random Forest es el más conservador (5 de 10 detectadas, 5 falsos negativos), Regresión Logística mejora a 6 detecciones (4 falsos negativos) y XGBoost se destaca con 7 detecciones correctas y solo 3 falsos negativos, lo que lo convierte en el modelo de mayor sensibilidad. Para un sistema de alerta temprana, esa reducción de falsos negativos es la diferencia más relevante, pues cada uno representa un atleta en riesgo que no recibiría intervención.

B.1.4. Comparación gráfica de métricas

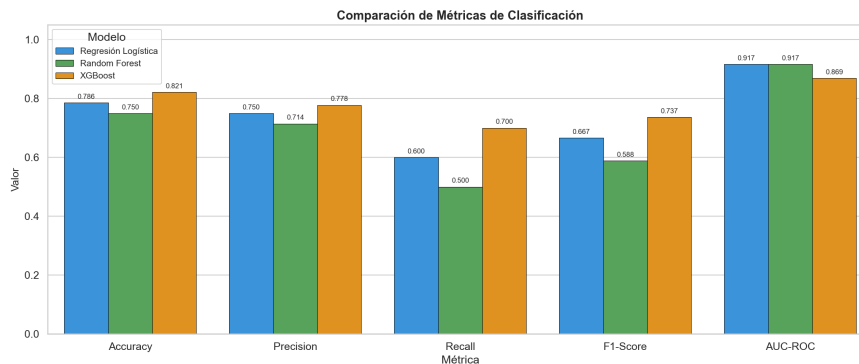


Figura B.2: Comparación de las cinco métricas de clasificación para los tres modelos.

La visualización de barras agrupadas confirma el patrón de la tabla comparativa. XGBoost domina en cuatro de las cinco métricas (*accuracy*, *precision*, *recall* y F1-Score). La diferencia más pronunciada se observa en el *recall*, donde la brecha entre Random Forest (0.500) y XGBoost (0.700) es de 20 puntos porcentuales, lo que en la práctica significa detectar 2 nadadores adicionales en riesgo de cada 10. Únicamente en AUC-ROC los otros dos modelos obtienen una ventaja marginal (0.917 frente a 0.869), relacionada con la calibración de probabilidades más que con la efectividad de las predicciones finales.

B.1.5. Curvas ROC

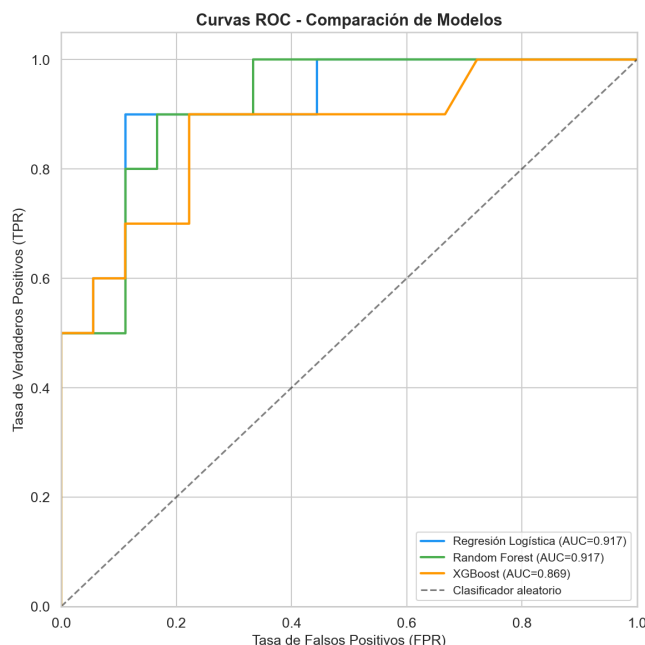


Figura B.3: Curvas ROC de los tres modelos comparadas con el clasificador aleatorio.

La curva ROC (*Receiver Operating Characteristic*) representa el compromiso entre la tasa de verdaderos positivos (proporción de nadadores en riesgo correctamente detectados) y la tasa de falsos positivos (proporción de nadadores normales clasificados erróneamente como en riesgo) a lo largo de distintos umbrales de decisión. La línea diagonal gris punteada representa un clasificador aleatorio; cualquier modelo útil debe situar su curva por encima. Las tres curvas se ubican significativamente por encima, confirmando capacidad real de discriminación. Las curvas de Regresión Logística y Random Forest ascienden de manera más pronunciada hacia la esquina superior izquierda, lo que se traduce en un AUC ligeramente mayor (0.917). La curva de XGBoost, con un AUC de 0.869, refleja que para ciertos rangos de umbral logra una alta tasa de detección con pocas falsas alarmas, lo que explica su superioridad con el umbral estándar de 0.5.

B.1.6. Curvas Precision-Recall

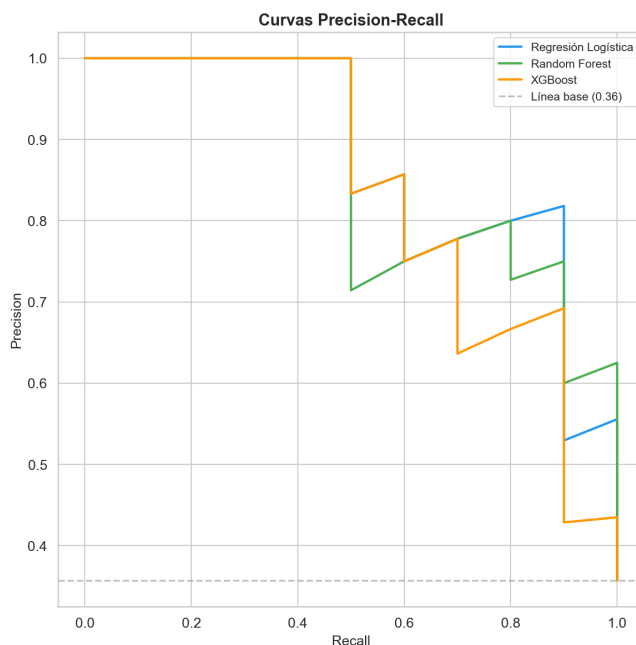


Figura B.4: Curvas Precision-Recall para los tres modelos.

Las curvas Precision-Recall resultan especialmente informativas con clases desbalanceadas, como aquí, donde el 36 % de las sesiones corresponden a bajo rendimiento. La línea base gris punteada (0.36) representa un clasificador trivial que siempre predice “bajo rendimiento”: lograría un *recall* de 1.0 pero una precisión de solo 0.36. El aspecto más destacable es el comportamiento de XGBoost, que mantiene una precisión perfecta (1.0) hasta un *recall* de aproximadamente 0.50, lo que significa que sus predicciones de mayor confianza no producen falsas alarmas; a partir de ese punto, al intentar detectar más nadadores en riesgo, la precisión desciende gradualmente. Este patrón habilita una decisión de diseño concreta para la plataforma: una política de alertas de dos niveles, en la que las predicciones de mayor certeza (el tramo de precisión máxima) se traten como alertas de alta confianza que justifican intervención directa, mientras que las de confianza moderada se presenten como señales para revisar, sujetas al juicio del entrenador.

B.1.7. Análisis de importancia de variables

La identificación de las variables que más influyen en cada modelo ofrece información potencialmente accionable para los entrenadores. Su lectura, no obstante, debe hacerse a la luz de una advertencia metodológica central, ya señalada en el Capítulo 4: varias de las variables más influyentes forman parte de la definición de la propia variable objetivo, por lo que su importancia está condicionada por una fuga de información. Los tres modelos proporcionan mecanismos distintos

para cuantificar esta importancia; se describen a continuación y se interpretan en conjunto al final de la subsección.

Importancia de variables en Random Forest. Random Forest estima la importancia de cada variable midiendo cuánto contribuye, en promedio, a la reducción de la impureza (índice de Gini) en las divisiones de todos los árboles del bosque.

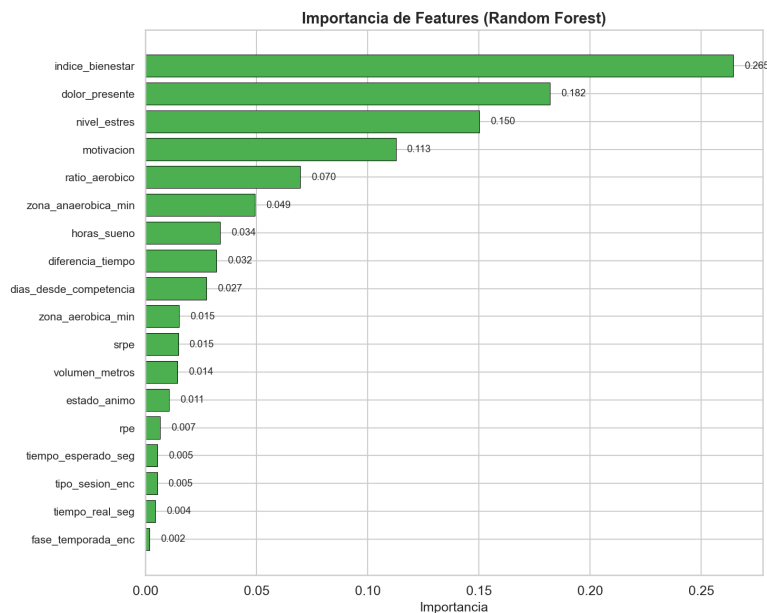


Figura B.5: Importancia de variables según Random Forest.

Los resultados revelan que el **índice de bienestar** es la variable más importante (0.265), seguida de la **presencia de dolor** (0.182), el **nivel de estrés** (0.150) y la **motivación** (0.113). Estas cuatro variables concentran más del 71 % de la capacidad predictiva del modelo. En contraste, variables de carga física como el volumen nadado (0.014), los tiempos cronométricos (< 0.005) y el tipo de sesión (0.005) muestran una importancia marginal.

Coefficientes de la Regresión Logística. La Regresión Logística ofrece una interpretación directa de la dirección del efecto de cada variable. Los coeficientes positivos (en rojo) indican que la variable aumenta la probabilidad de bajo rendimiento, mientras que los negativos (en azul) la reducen.

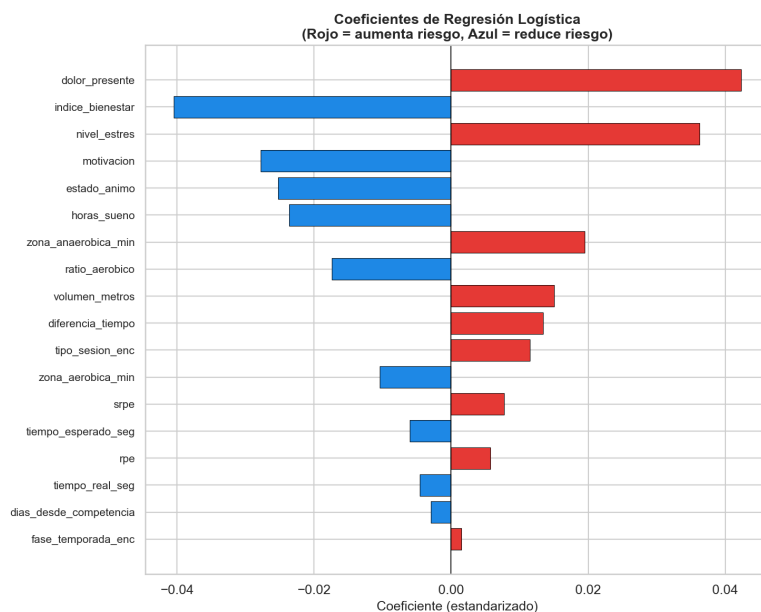


Figura B.6: Coeficientes estandarizados de la Regresión Logística.

La presencia de dolor es el mayor factor de riesgo (coeficiente positivo más alto), seguido del nivel de estrés y los minutos en zona anaeróbica. Como factores protectores destacan el índice de bienestar (coeficiente negativo más alto), la motivación, el estado de ánimo positivo y las horas de sueño. La *dirección* de estos efectos es coherente con la literatura sobre sobreentrenamiento y fatiga deportiva, que identifica los factores psicológicos y de bienestar como predictores tempranos de deterioro del rendimiento [Bernaciková et al., 2022](#); [Barry et al., 2024](#); su *magnitud*, en cambio, está condicionada por la fuga que se discute al cierre de la subsección.

Importancia de variables en XGBoost. XGBoost mide la importancia de cada variable a través de la ganancia (*gain*), que cuantifica la mejora promedio en la función de pérdida cuando la variable se utiliza para realizar una división en los árboles.

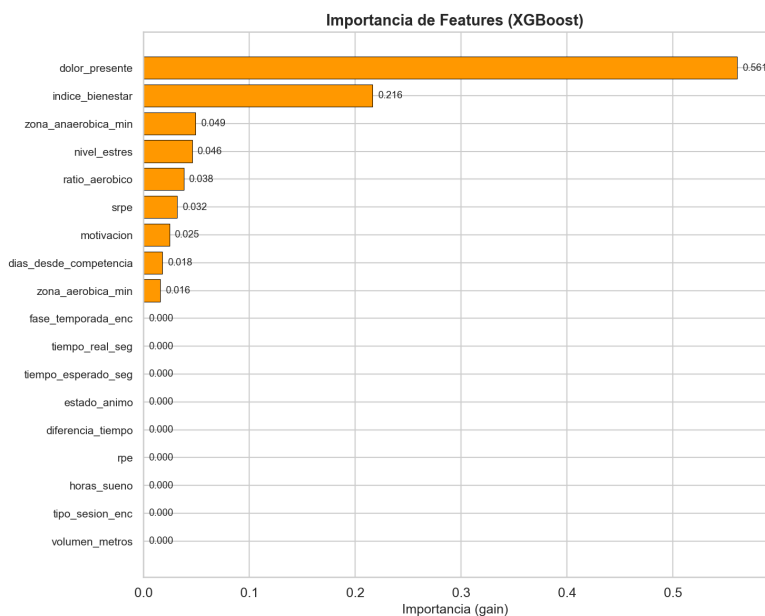


Figura B.7: Importancia de variables según XGBoost (ganancia).

Los resultados de XGBoost presentan una concentración notable: la **presencia de dolor** domina con una ganancia de 0.561 (56.1 % de la importancia total), seguida del **índice de bienestar** con 0.216 (21.6 %). Juntas, estas dos variables representan el 77.7 % de la capacidad predictiva del modelo. Las demás contribuyen marginalmente, y 10 de las 18 variables reciben una importancia de 0.000, lo que indica que XGBoost las considera redundantes una vez que dispone de la información de dolor y bienestar. Bajo el esquema actual, los factores más asociados a la predicción no son las métricas de carga física (volumen, zonas de entrenamiento), sino las variables de **estado del nadador** (dolor, bienestar, estrés, motivación). Ahora bien, como varias de esas variables intervienen en la definición de la etiqueta, este resultado no puede leerse todavía como evidencia de que el estado del nadador determine empíricamente el bajo rendimiento; debe tratarse como una hipótesis a verificar. Lo que sí es consistente es el énfasis del instrumento en el registro diario del bienestar, cuyo valor predictivo real deberá confirmarse sin la fuga descrita.

Lectura de la importancia bajo la fuga de información. Los tres modelos coinciden en señalar el dolor, el bienestar, el estrés y la motivación como las variables más influyentes. Esta coincidencia es esperable y, en buena medida, circular: esas mismas variables intervienen en la construcción de la variable objetivo (los siete criterios de riesgo del Capítulo 4), de modo que el modelo dispone, entre sus predictoras, de información que define la etiqueta y puede en parte reconstruirla en lugar de predecirla. En consecuencia, tanto la magnitud de las importancias como el nivel absoluto de las métricas reportadas, en validación cruzada y en prueba, deben considerarse optimistas. La dominancia observada del estado del nadador es, por tanto, una hipótesis a verificar sobre un corpus real más amplio y excluyendo las variables constitutivas del objetivo, no una con-

clusión consolidada. La Fase 2 ilustra esta corrección: excluye explícitamente `diferencia_tiempo`, `srpe` e `indice_bienestar` (Cuadro B.5), y el mismo criterio guiará la recalibración de la Fase 1 cuando se disponga de un corpus real más amplio.

B.1.8. Análisis de sobreajuste

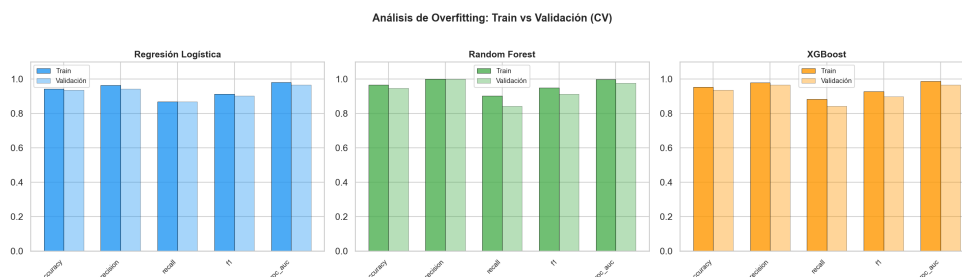


Figura B.8: Análisis de *overfitting*: comparación de métricas en entrenamiento versus validación cruzada.

El sobreajuste ocurre cuando un modelo aprende patrones específicos del entrenamiento que no se generalizan a datos nuevos. Para evaluarlo se comparan las métricas de entrenamiento (barras opacas) con las de validación cruzada (barras transparentes); un sobreajuste severo se manifiesta como una brecha pronunciada entre ambas.

Regresión Logística presenta la menor brecha en todas las métricas, lo que indica una excelente capacidad de generalización y un sobreajuste mínimo, consistente con su naturaleza lineal y su capacidad limitada para memorizar patrones complejos. **Random Forest** muestra la mayor brecha, particularmente en *precision* y *recall*, donde el entrenamiento alcanza valores cercanos a 1.0 frente a una validación en torno a 0.85 a 0.95, sugiriendo un sobreajuste moderado esperable en modelos basados en árboles. **XGBoost** presenta un comportamiento intermedio, con brechas controladas gracias a sus mecanismos de regularización (`reg_alpha`, `reg_lambda`, `subsample`, `early stopping`), manteniendo entrenamiento y validación razonablemente alineados.

Conviene señalar que los tres modelos muestran métricas de validación cruzada (F1 medio ≈ 0.91) superiores a las del conjunto de prueba final (F1 entre 0.59 y 0.74). Esta diferencia es esperable y no invalida los resultados: el conjunto de prueba contiene únicamente 28 observaciones, por lo que la variabilidad estadística es inherentemente mayor; con más datos en el futuro se espera una convergencia más estrecha entre ambas estimaciones. Debe distinguirse, además, este efecto de tamaño muestral del señalado en la subsección de importancia: la fuga de información eleva el nivel *absoluto* de todas estas métricas (tanto en validación cruzada como en prueba), mientras que la brecha entre ambas responde principalmente al reducido tamaño del conjunto de prueba.

B.2. Fase 2: detalle del modelo MLP de predicción del rendimiento futuro

B.2.1. Pipeline de preprocesamiento del dataset sintético

El preprocesamiento transformó el dataset sintético en una base limpia, estructurada y apta para modelado supervisado de regresión. Partiendo de 1 680 registros y 17 columnas, se eliminaron 121 filas correspondientes a sesiones de inasistencia, dado que no contenían datos de rendimiento y su imputación habría introducido ruido artificial; se verificó además la ausencia de duplicados exactos y de inconsistencias lógicas como tiempos negativos o valores fisiológicos fuera de rango. Tras la depuración y la construcción de variables temporales necesarias para predecir $t + 1$, se conservaron 1 519 filas útiles. La partición se realizó mediante un *split* temporal estricto basado en fechas globales (70 % entrenamiento, 15 % validación, 15 % prueba), y la estandarización con **StandardScaler** se ajustó exclusivamente sobre el conjunto de entrenamiento, evitando la fuga de estadísticas futuras hacia el aprendizaje [Bischi et al., 2023](#).

Tabla B.4: Decisiones del pipeline de preprocesamiento de la Fase 2.

Etapas	Decisión aplicada	Justificación metodológica
Inasistencias	Eliminación de 121 filas sin datos de sesión.	Evita introducir ruido por imputación de registros que no representan rendimiento.
Duplicados	Verificación de ausencia de duplicados exactos (0 detectados).	Previene la sobrerepresentación accidental de observaciones idénticas.
Validación de rangos	Verificación de tiempos no negativos y rangos fisiológicos.	Garantiza integridad lógica antes del modelado.
Columnas problemáticas	Exclusión de <code>diferencia_tiempo</code> , <code>srpe</code> e <code>indice_bienestar</code> .	Elimina fuga de información y redundancia con otras variables.
División temporal	70 % entrenamiento, 15 % validación, 15 % prueba por fecha global.	Preserva el orden cronológico y emula condiciones reales de uso.
Estandarización	StandardScaler ajustado solo sobre entrenamiento.	Evita fuga de estadísticas del futuro hacia el aprendizaje.

Tabla B.5: Columnas excluidas del modelado y motivo de exclusión.

Columna excluida	Motivo de exclusión
diferencia_tiempo	Fuga de información: se calcula como <code>tiempo_real - tiempo_plan</code> , por lo que contiene el target.
srpe	Redundancia con <code>rpe</code> y <code>volumen_total</code> , ya presentes como predictores.
indice_bienestar	Redundancia con sus componentes base: dolor, estado de ánimo, motivación y estrés.

Estas decisiones buscan preservar la validez predictiva del modelo, evitando que aprenda relaciones artificiales o información que no estaría disponible en un escenario real de uso. Es importante notar el contraste con la Fase 1: mientras allí las variables constitutivas del objetivo permanecieron entre las predictoras (lo que motivó la lectura cautelosa de su importancia), aquí se excluyen de forma explícita, anticipando el tratamiento que recibirá la Fase 1 en su recalibración. Conviene recordar, en todo caso, que el dataset sigue siendo sintético; el rigor del preprocesamiento no compensa esa naturaleza, pero garantiza que la evaluación posterior refleje la capacidad del modelo y no artefactos del proceso.

B.2.2. Selección de variables predictoras

La selección de variables se orientó a incluir únicamente información disponible al cierre del día actual t , con el propósito de predecir el rendimiento de la sesión siguiente $t + 1$. Se definieron variables numéricas de rendimiento, percepción y carga; variables temporales derivadas que capturan tendencias recientes; y variables categóricas codificadas con `LabelEncoder`.

Tabla B.6: Variables predictoras seleccionadas para el modelo MLP.

Variable	Tipo	Justificación
tiempo_100m_real_sec	Numérica	Rendimiento observado en t ; antecedente directo del rendimiento futuro.
tiempo_100m_plan_sec	Numérica	Objetivo planificado por el entrenador; representa el nivel esperado.
rpe	Numérica	Esfuerzo percibido durante la sesión actual.
horas_sueno	Numérica	Indicador de recuperación.
nivel_estres	Numérica	Estado emocional del día.
motivacion	Numérica	Disposición psicológica para entrenar.
volumen_total	Numérica	Carga externa de entrenamiento.
lag_1_tiempo	Temporal	Tiempo real de la sesión anterior.
rolling_mean_3	Temporal	Media móvil de los últimos tres tiempos.
rolling_std_3	Temporal	Variabilidad reciente del rendimiento.
rolling_mean_rpe_3	Temporal	Carga perceptual reciente acumulada.
rolling_mean_sueno_3	Temporal	Calidad reciente del descanso.
nadador_nombre	Categórica	Identifica al nadador; cada uno posee un nivel base distinto.
estilo	Categórica	Determina el rango de tiempos esperable.
tipo_sesion	Categórica	Tipo de entrenamiento (R1, R2, VO2, entre otros).
dolor	Categórica binaria	Presencia de dolor; puede afectar el rendimiento.
estado_animo	Categórica ordinal	Estado emocional del nadador.

Las variables categóricas se codificaron porque aportan información estructural relevante: cada nadador posee un nivel base distinto, cada estilo tiene rangos de tiempo específicos y las condiciones de dolor o estado emocional pueden modificar la respuesta al entrenamiento. La codificación ordinal mediante `LabelEncoder` fue suficiente para este caso, dado que el modelo aprende a tratar estas variables como factores discretos sin imponer una métrica continua sobre ellas.

Observación metodológica sobre la correlación entre plan y target. Se identificó una correlación de 0.97 entre `tiempo_100m_plan_sec` y el target. Esta correlación debe interpretarse de manera doble: como fortaleza, porque el tiempo planificado contiene información altamente predictiva sobre el rendimiento esperado; y como limitación, porque su dominancia puede reducir la visibilidad del aporte de variables de bienestar o carga durante el aprendizaje. Esta tensión se

discute en la sección de limitaciones de la Fase 2.

B.2.3. Comparación de alternativas y arquitectura del modelo

Se seleccionó un *Multi-Layer Perceptron* (MLP) denso como modelo de regresión, tras evaluar alternativas recurrentes (LSTM, GRU) que fueron descartadas explícitamente. Las arquitecturas recurrentes resultaron poco adecuadas porque el dataset contiene únicamente 1 519 filas útiles distribuidas en 40 combinaciones de nadador y estilo, con entre 10 y 65 registros cada una, volumen insuficiente para entrenar secuencias recurrentes sin un riesgo elevado de memorización. Además, los *lags* explícitos y las ventanas móviles ya codifican la dependencia temporal relevante, lo que permite a un modelo tabular capturar dicha información sin reconstruir secuencias completas; finalmente, el MLP trabaja con un vector plano de variables, lo que simplifica la inferencia en producción y favorece la integración en la plataforma web modular.

Tabla B.7: Comparación de alternativas evaluadas para el modelo de regresión.

Alternativa de modelo	Justificación
LSTM / GRU	No seleccionados. La escasez de secuencias estables por combinación de nadador y estilo eleva el riesgo de sobreajuste y de memorización de trayectorias individuales.
MLP tabular	Seleccionado. Compatible con <i>features</i> tabulares (incluidos lags y ventanas móviles), de menor complejidad y con despliegue más sencillo en la plataforma.

La arquitectura final está compuesta por tres capas ocultas densas de 64, 32 y 16 neuronas con activación ReLU. Para controlar la complejidad y mitigar el sobreajuste, se incorporaron regularización L2 sobre los pesos y *dropout* entre capas. La capa de salida es lineal y de una sola neurona, coherente con la tarea de regresión sobre el tiempo real en segundos.

Tabla B.8: Componentes principales de la arquitectura y entrenamiento del modelo MLP.

Componente	Configuración	Función dentro del modelo
Capa oculta 1	64 neuronas, ReLU	Captura interacciones de alto nivel entre variables.
Capa oculta 2	32 neuronas, ReLU	Reduce dimensionalidad y refina representaciones intermedias.
Capa oculta 3	16 neuronas, ReLU	Concentra la información antes de la salida.
Regularización L2	Aplicada sobre los pesos	Penaliza pesos grandes y reduce sobreajuste.
Dropout	Aplicado entre capas	Introduce robustez al desactivar neuronas durante el entrenamiento.
Capa de salida	1 neurona lineal	Predice el tiempo real en segundos.
Función de pérdida	MAE	Error medio absoluto, interpretable en segundos.
Optimizador	Adam, lr inicial = 0,001	Adapta la tasa de aprendizaje durante el entrenamiento.
EarlyStopping	Activo	Detiene el entrenamiento al estancarse la métrica de validación.
ReduceLROnPlateau	Activo	Reduce la tasa de aprendizaje cuando la métrica deja de mejorar.
ModelCheckpoint	Activo	Conserva los pesos del mejor modelo durante el entrenamiento.

El uso de MAE como función de pérdida facilita la interpretación del error en segundos, una unidad directamente comprensible para los entrenadores. Esta decisión resulta especialmente útil en el contexto deportivo, ya que un error de pocos segundos en la predicción puede traducirse a una magnitud que el cuerpo técnico evalúa con criterios profesionales, sin necesidad de transformaciones adicionales [Chicco and Jurman, 2020](#).

B.2.4. Interpretación gráfica de los resultados

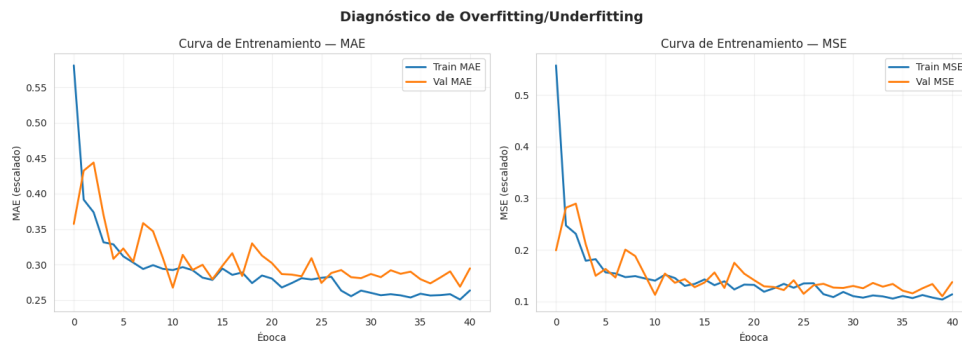


Figura B.9: Curvas de entrenamiento y validación del modelo MLP para MAE y MSE.

Las curvas de entrenamiento muestran una reducción rápida del error durante las primeras épocas, seguida de una etapa de estabilización. La cercanía entre las trayectorias de entrenamiento y validación indica que el modelo aprende los patrones dominantes sin presentar una divergencia creciente entre ambos conjuntos. Aunque existe una oscilación moderada en la validación, esta es consistente con la variabilidad natural entre nadadores y estilos, y sugiere que la regularización y el *early stopping* limitaron adecuadamente la complejidad del modelo.

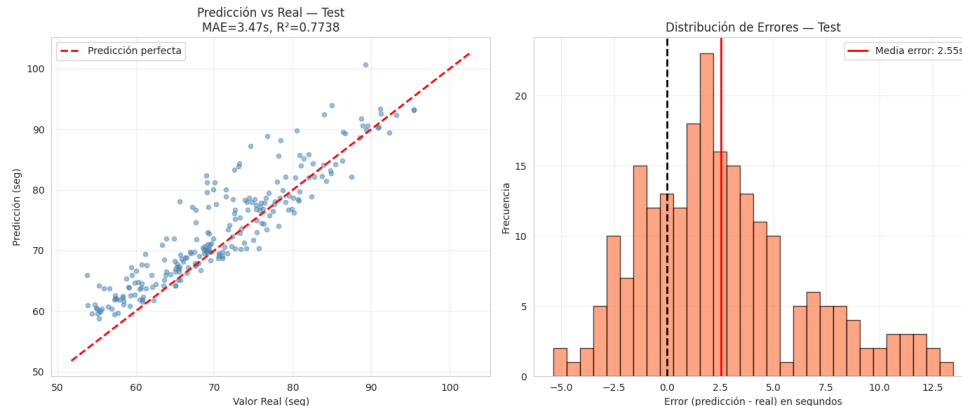


Figura B.10: Relación entre valores reales y predichos, y distribución de errores en el conjunto de prueba.

El diagrama de predicción frente a valor real evidencia una alineación general con la diagonal de referencia, lo que indica que el modelo captura razonablemente la magnitud del rendimiento futuro. No obstante, la dispersión aumenta en los tiempos más altos, lo que sugiere menor precisión en escenarios menos frecuentes dentro del dataset. Por su parte, el histograma de errores muestra una media del error, calculada como valor predicho menos valor real, de aproximadamente +2,55 s; es decir, el modelo tiende a sobreestimar el tiempo y, por tanto, a pronosticar en promedio un rendimiento ligeramente peor del que ocurre. En términos operativos este sesgo es *conservador* para

un sistema de alerta, porque inclina las predicciones hacia el lado de la precaución; sin embargo, un sesgo sistemático de esta magnitud podría inducir fatiga de alarmas si la plataforma reacciona ante cada sobreestimación, por lo que su corrección es uno de los objetivos de la recalibración con datos reales. Aun así, la concentración central de los errores indica que la predicción se mantiene en rangos operativamente manejables para apoyar decisiones de corto plazo.

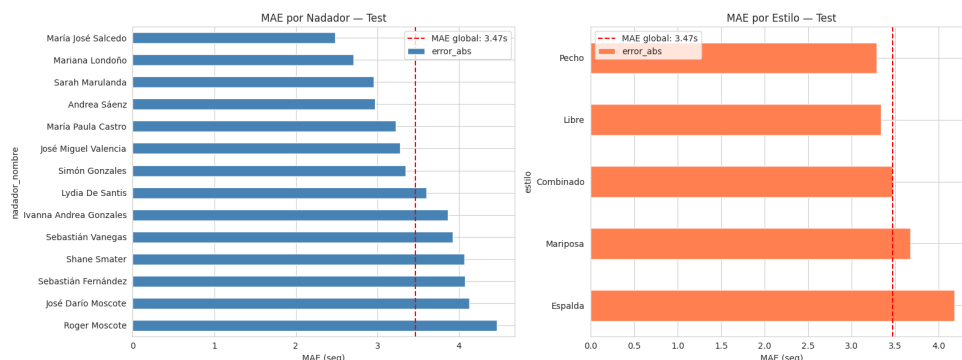


Figura B.11: Error absoluto medio del modelo MLP por nadador y por estilo en el conjunto de prueba.

El análisis por nadador muestra una variabilidad importante en el MAE, lo cual refleja diferencias en la estabilidad intraindividual y en la cantidad de historial disponible para cada perfil. Las predicciones tienden a ser más precisas en nadadores con trayectorias regulares y menos precisas en aquellos con oscilaciones pronunciadas; el rango observado va desde el nadador con menor MAE hasta el nadador con mayor MAE dentro del conjunto evaluado. Por estilo, los errores se mantienen en un rango más homogéneo, aproximadamente entre 3.31 y 4.07 segundos, lo que sugiere que el modelo no favorece de manera marcada una técnica específica y generaliza razonablemente entre estilos de nado. Estas diferencias deben leerse como propiedades del comportamiento predictivo del modelo y no como una evaluación personal de los atletas.

Metodología de entrenamiento utilizada en Fabio Toro Academy

C.1. Propósito del anexo

El presente anexo complementa la sección de rendimiento deportivo y entrenamiento por zonas en natación, incorporando la metodología de planificación utilizada en Fabio Toro Academy durante el periodo considerado para el desarrollo del trabajo de grado. Su propósito no es formular un modelo universal de entrenamiento en natación carreras, sino describir la lógica operativa del caso de estudio, de manera que se pueda comprender cómo se estructuraban las cargas, las zonas de trabajo, la frecuencia semanal y la programación del macrociclo que sirvió como referencia para la caracterización de datos y el diseño del componente analítico de la plataforma.

C.2. Organización general del entrenamiento

Durante el periodo analizado, la planificación de Fabio Toro Academy combinó trabajo en agua, y seguimiento del estado del nadador. En condiciones regulares de carga, la dinámica semanal contempló dos sesiones diarias de lunes a viernes y una sesión diaria durante sábado y domingo, lo que corresponde a una referencia de doce sesiones semanales.

Tabla C.1: Componentes generales de la metodología de entrenamiento en Fabio Toro Academy

Componente	Descripción operativa	Relación con el seguimiento del rendimiento
Trabajo en agua	Comprende las sesiones principales de natación, organizadas por zonas, capacidades, volúmenes semanales, ejercicios técnicos, series principales, patada, recuperación y uso de implementos.	Permite registrar variables de carga externa como volumen, zona de entrenamiento, tipo de sesión, capacidad trabajada, tiempo planificado y tiempo real.
Seguimiento subjetivo	Incluye variables reportadas por el nadador, como percepción del esfuerzo, estado de ánimo, motivación, estrés, sueño y presencia de dolor.	Aporta información sobre carga interna y estado del deportista, relevante para anticipar cambios de rendimiento y orientar alertas tempranas.
Control y competencia	Incluye semanas o sesiones destinadas a pruebas de control, ajuste de cargas, preparación competitiva y participación en eventos del calendario deportivo.	Permite relacionar el comportamiento del nadador con fases específicas del macrociclo y con la proximidad a eventos competitivos.

C.3. Zonas y capacidades de entrenamiento

La planificación implementada por la academia utiliza una nomenclatura basada en capacidades de entrenamiento. Estas capacidades se expresan mediante abreviaturas que permiten identificar el propósito dominante de cada semana o sesión. Aunque una sesión puede contener bloques complementarios de técnica, velocidad o recuperación, la capacidad asignada indica el énfasis principal del trabajo.

Tabla C.2: Nomenclatura de zonas y capacidades de entrenamiento

Nomenclatura	Nombre operativo	Significado dentro de la planificación
R1	Resistencia aeróbica básica	Zona de base aeróbica. Se utiliza para acumular volumen, mejorar economía de nado, sostener ritmo controlado y reforzar aspectos técnicos bajo intensidad moderada.
R1-2	Transición aeróbica	Zona intermedia entre el trabajo aeróbico básico y un estímulo aeróbico de mayor exigencia. Permite progresar la carga sin pasar directamente a intensidades altas.
R2	Resistencia aeróbica de desarrollo	Trabajo aeróbico más exigente que R1. Busca incrementar la capacidad de sostener ritmos más demandantes, manteniendo control técnico y tolerancia al volumen.
MVO2 / VO2	Consumo máximo de oxígeno	Trabajo de alta exigencia aeróbica. Se utiliza para estimular la potencia aeróbica, mejorar la capacidad de sostener ritmos elevados y preparar al nadador para demandas competitivas.

Continúa en la siguiente página

Continuación de la Tabla C.2

Nomenclatura	Nombre operativo	Significado dentro de la planificación
RL	Resistencia láctáida	Trabajo orientado a sostener esfuerzos intensos con acumulación progresiva de fatiga. Se emplea para preparar al nadador ante tramos de carrera donde debe mantener velocidad bajo estrés metabólico.
TL	Tolerancia láctáida	Trabajo de alta intensidad asociado a la capacidad de tolerar esfuerzos exigentes y sostener la técnica en condiciones de fatiga. Su presencia suele ubicarse cerca de fases específicas o competitivas.

Esta nomenclatura fue relevante para el trabajo de grado porque permitió traducir la planificación deportiva a variables estructuradas dentro del sistema. En particular, las zonas de entrenamiento aportaron contexto para interpretar la carga externa, mientras que las variables de bienestar permitieron complementar dicha lectura con información sobre el estado del nadador.

C.4. Programación del macrociclo diciembre–marzo

La programación entregada por Fabio Toro Academy corresponde a un macrociclo organizado en diecisiete semanas, identificado en la planificación como semanas 24 a 40. El bloque inicia en diciembre y se extiende hasta la primera semana de abril. La planificación contempla tres momentos generales: una fase inicial de transición y adaptación, una fase de incremento y sostenimiento de la carga, y una fase final orientada a estímulos específicos, competencia y transición posterior.

Tabla C.3: Macro ciclo utilizado como referencia en el caso de estudio

Periodo	Semanas	Fechas aproximadas	Capacidades predominantes	Interpretación metodológica
Diciembre	24–27	8 de diciembre – 4 de enero	R1, R1-2	Fase de transición, regeneración y readaptación. Se manejan volúmenes reducidos cercanos a 21 km semanales y una frecuencia de seis sesiones por semana.
Enero	28–31	5 de enero – 1 de febrero	R2, VO2	Fase de incremento de carga. Se pasa a una referencia de doce sesiones semanales, con volúmenes entre 60 y 70 km, orientados al desarrollo aeróbico y a estímulos de potencia aeróbica.
Febrero	32–35	2 de febrero – 1 de marzo	RL, R2	Fase de sostenimiento y verificación. Se combinan cargas aeróbicas, estímulos lácticos y una semana de control para revisar la respuesta del nadador.
Marzo	36–39	2 de marzo – 29 de marzo	VO2, RL, TL	Fase específica y competitiva. Se incorporan estímulos de VO2, resistencia láctica y tolerancia láctica, con reducción progresiva del volumen antes de la competencia Ciudad de Cali.
Cierre	40	30 de marzo – 5 de abril	R1	Fase posterior a la competencia, orientada a transición y recuperación antes de iniciar una nueva planificación.

En términos de volumen, la programación reporta un total de 795 km de trabajo en agua para el macro ciclo. La mayor proporción corresponde a R1, con 522,3 km, equivalente al 65,7 % del volumen total.

C.5. Relación de la metodología de entrenamiento con el prototipo desarrollado 123

Tabla C.4: Distribución del volumen en agua por zona durante el macrociclo

Zona	Volumen total (km)	Participación	Lectura metodológica
R1	522,3	65,7 %	Base principal del proceso. Prioriza volumen, continuidad, eficiencia técnica y control del ritmo.
R1-2	64,0	8,1 %	Zona de transición aeróbica para progresar la carga sin elevarla de forma abrupta.
R2	66,5	8,4 %	Desarrollo aeróbico de mayor exigencia, utilizado para sostener ritmos más demandantes.
MVO2	42,0	5,3 %	Estímulo de potencia aeróbica, ubicado en semanas específicas de carga.
RL	21,0	2,6 %	Trabajo láctico de resistencia, empleado de forma controlada.
TL	11,0	1,4 %	Trabajo de tolerancia láctica, concentrado cerca de la fase competitiva.
Velocidad	8,2	1,0 %	Estímulos breves de calidad, orientados a potencia, frecuencia y eficiencia.
Regeneración	60,0	7,5 %	Recuperación activa y asimilación de cargas durante el ciclo.
Total	795,0	100 %	Volumen total programado en agua durante el macrociclo.

La distribución confirma una planificación predominantemente aeróbica, coherente con un proceso formativo de nadadores juveniles en el que se busca construir base, sostener adaptación progresiva y reservar los estímulos de mayor intensidad para momentos específicos del ciclo. Esta lectura también permite comprender por qué, dentro del sistema desarrollado, el volumen y las zonas de entrenamiento debían registrarse junto con variables subjetivas de bienestar, ya que la carga externa por sí sola no explica completamente el estado del nadador.

C.5. Relación de la metodología de entrenamiento con el prototipo desarrollado

La metodología de entrenamiento de Fabio Toro Academy sirvió como base contextual para definir las variables relevantes del sistema y para comprender la forma en que el cuerpo técnico interpreta el rendimiento.

En consecuencia, la plataforma desarrollada no se orientó solo a almacenar resultados de tiempos, sino a centralizar información de entrenamiento, bienestar y seguimiento. Al incorporar la

metodología de entrenamiento como contexto, el sistema puede representar de forma más clara la relación entre carga planificada, ejecución de la sesión, estado del nadador y posibles señales de bajo rendimiento.