

BIOESTADÍSTICA APLICADA EN LA ENSEÑANZA DE LAS CIENCIAS DE LA SALUD



AUTORES:

Castro Jalca Jazmín Elena
Azúa Menéndez Marieta Del Jesús
Jijón Cañarte Lidia Fernanda
Hidalgo Villavicencio Gilson Alfonso

Chilán Santana Caleb Isaac
Mina Ortiz Jhon Brian
Regalado Jalca Julio Johnny





Primera Edición 2025

ISBN: 978-9942-571-13-7

2025, RUNAIKI Editora-Editorial Internacional. España.

<https://runaiki.es/index.php/runaiki>

Consejo Editorial:

Dra. Mayra Díaz Martínez, PhD.

RUNAiki Editora-Editorial Internacional. España.

Hecho en Ecuador

Este texto ha sido sometido a un riguroso proceso de evaluación por pares externos.

Advertencia: Todos los derechos reservados. Queda prohibida la reproducción, registro o transmisión parcial o total de esta obra por cualquier sistema de recuperación de información existente o futuro, sin el permiso previo y por escrito del titular de los derechos correspondientes.

**Escanea el QR
para acceder al libro**



ISBN: 978-9942-571-13-7



Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

Autores Investigadores:

Castro Jalca Jazmín Elena

Licenciada en Laboratorio Clínico

Doctora en Ciencias de la Salud

Profesor titular de la Universidad Estatal del sur de Manabí

Jipijapa; Ecuador

 jazmin.castro@unesum.edu.ec

 <https://orcid.org/0000-0001-7593-8552>

Azúa Menéndez Marieta Del Jesús

Ingeniera en Computación y Redes

Magister en Educación Informática

Profesor de la Universidad Estatal del sur de Manabí

Jipijapa; Ecuador

 marieta.azua@unesum.edu.ec

 <https://orcid.org/0000-0002-5601-6621>

Jijón Cañarte Lidia Fernanda

Licenciada en Laboratorio Clínico

Magister en Ciencias del Laboratorio Clínico

Profesor de la Universidad Estatal del sur de Manabí

Jipijapa; Ecuador

 lidia.jijon@unesum.edu.ec

 <https://orcid.org/0000-0002-7780-5865>

Hidalgo Villavicencio Gilson Alfonso

Licenciado en Laboratorio Clínico

Magister en Ciencias del Laboratorio Clínico

Profesor de la Universidad Estatal del sur de Manabí

Jipijapa; Ecuador

 gilson.hidalgo@unesum.edu.ec

 <https://orcid.org/0000-0003-3364-7700>

Chilán Santana Caleb Isaac

Licenciado en Laboratorio Clínico

Magister en Biomedicina

Profesor de la Universidad Estatal del sur de Manabí

Jipijapa; Ecuador

✉ caleb.chilan@unesum.edu.ec

 <https://orcid.org/0000-0002-2832-8759>

Mina Ortiz Jhon Brian

Licenciado en Laboratorio Clínico

Magister en Biotecnología

Profesor titular de la Universidad Estatal del sur de Manabí

Jipijapa; Ecuador

✉ jhon.mina@unesum.edu.ec

 <https://orcid.org/0000-0002-3455-2503>

Regalado Jalca Julio Johnny

Ingeniero en Computación y Redes

Magister en Educación Informática

Profesor de la Universidad Estatal del sur de Manabí

Jipijapa; Ecuador

✉ julio.regalado@unesum.edu.ec

 <https://orcid.org/0000-0003-1669-8522>



Resumen

La bioestadística es una disciplina esencial para la comprensión, el análisis y la interpretación de datos en las ciencias de la salud, al permitir fundamentar decisiones clínicas, epidemiológicas y de laboratorio sobre bases cuantitativas rigurosas. Este libro presenta un tratamiento sistemático y aplicado de la bioestadística, integrando fundamentos conceptuales con procedimientos metodológicos orientados a la práctica sanitaria y a la formación universitaria. Se abordan los principios básicos del razonamiento estadístico, la organización y descripción de datos biomédicos, el uso de la probabilidad y de las distribuciones estadísticas más frecuentes en salud, así como los métodos de inferencia estadística empleados en la investigación clínica. El texto incorpora el análisis bioestadístico aplicado al laboratorio clínico, incluyendo el estudio de la variabilidad biológica y analítica, el control de calidad, la validación de métodos y la evaluación del desempeño diagnóstico mediante indicadores cuantitativos. Asimismo, se desarrolla el análisis epidemiológico desde una perspectiva poblacional, considerando medidas de frecuencia, asociación y síntesis de evidencia. Finalmente, se introduce el uso de inteligencia artificial y aprendizaje automático como herramientas complementarias para el análisis predictivo de datos clínicos, integradas dentro de un modelo algorítmico coherente con los principios estadísticos clásicos. El enfoque del libro prioriza la interpretación crítica de los resultados, la evaluación de la incertidumbre y la comunicación transparente de la evidencia, en concordancia con los principios de la medicina basada en evidencia. Esta obra está dirigida a estudiantes y profesionales de las ciencias de la salud y tiene como propósito fortalecer competencias analíticas y metodológicas que contribuyan a una práctica sanitaria informada, responsable y basada en datos.

Palabras clave: Bioestadística; ciencias de la salud; análisis de datos biomédicos; inferencia estadística; epidemiología; aprendizaje automático; medicina basada en evidencia.

Abstract

Biostatistics is an essential discipline for the understanding, analysis, and interpretation of data in the health sciences, as it provides a quantitative basis for clinical, epidemiological, and laboratory decision-making. This book presents a systematic and applied treatment of biostatistics, integrating conceptual foundations with methodological procedures oriented towards health practice and university education. It addresses the basic principles of statistical reasoning, the organisation and description of biomedical data, the use of probability and statistical distributions commonly applied in health research, and inferential statistical methods used in clinical investigation. The text incorporates applied biostatistics in the clinical laboratory, including the study of biological and analytical variability, quality control, method validation, and the evaluation of diagnostic performance through quantitative indicators. It also develops epidemiological analysis from a population-based perspective, considering measures of frequency, association, and evidence synthesis. In addition, the book introduces artificial intelligence and machine learning as complementary tools for predictive analysis of clinical data, integrated within an algorithmic framework consistent with classical statistical principles. Emphasis is placed on the critical interpretation of results, the assessment of uncertainty, and the transparent communication of evidence, in accordance with the principles of evidence-based medicine. This book is intended for students and professionals in the health sciences and aims to strengthen analytical and methodological competencies that support informed, responsible, and data-driven health practice.

Keywords: Biostatistics; health sciences; biomedical data analysis; statistical inference; epidemiology; machine learning; evidence-based medicine.

ÍNDICE

RESUMEN	IV
ABSTRACT.....	V
PRÓLOGO	IX
1 FUNDAMENTOS DE LA BIOESTADÍSTICA EN LAS CIENCIAS DE LA SALUD..	12
1.1 CONCEPTO Y EVOLUCIÓN HISTÓRICA.....	12
1.2 IMPORTANCIA DE LA BIOESTADÍSTICA EN LA INVESTIGACIÓN BIOMÉDICA Y CLÍNICA.....	18
1.3 TIPOS DE VARIABLES: CUALITATIVAS, CUANTITATIVAS Y SUS ESCALAS DE MEDICIÓN..	26
1.4 POBLACIÓN, MUESTRA Y PARÁMETROS ESTADÍSTICOS	33
1.5 FUENTES DE DATOS EN SALUD: CLÍNICOS, EPIDEMIOLOGICOS Y EXPERIMENTALES....	39
1.6 ERRORES DE MEDICIÓN Y SESGOS EN LOS ESTUDIOS DE SALUD	43
1.7 ÉTICA Y CONFIDENCIALIDAD EN EL MANEJO DE DATOS SANITARIOS.....	47
2 ORGANIZACIÓN Y DESCRIPCIÓN DE DATOS BIOMÉDICOS.....	50
2.1 TABULACIÓN Y CLASIFICACIÓN DE DATOS CLÍNICOS	50
2.2 TABLAS DE FRECUENCIA Y REPRESENTACIÓN DE RESULTADOS DE LABORATORIO.	56
2.3 DIAGRAMAS DE BARRAS, HISTOGRAMAS Y POLÍGONOS DE FRECUENCIA	62
2.4 DIAGRAMAS DE DISPERSIÓN Y CORRELACIÓN VISUAL	67
2.5 MEDIDAS DE TENDENCIA CENTRAL: MEDIA, MEDIANA Y MODA.....	72
2.6 MEDIDAS DE DISPERSIÓN: RANGO, VARIANZA Y DESVIACIÓN ESTÁNDAR.....	75
2.7 APLICACIONES EN SOFTWARE ESTADÍSTICO	79
3 PROBABILIDAD Y DISTRIBUCIONES EN SALUD	86
3.1 CONCEPTO DE PROBABILIDAD EN EL CONTEXTO SANITARIO	86
3.2 REGLA DE ADICIÓN Y MULTIPLICACIÓN DE PROBABILIDADES	90
3.3 TEOREMA DE BAYES APLICADO AL DIAGNÓSTICO CLÍNICO	94

3.4	DISTRIBUCIÓN BINOMIAL Y APLICACIÓN A RESULTADOS POSITIVOS Y NEGATIVOS.....	99
3.5	DISTRIBUCIÓN NORMAL Y SU RELEVANCIA EN VALORES DE REFERENCIA DE LABORATORIO	103
3.6	DISTRIBUCIONES T DE STUDENT, CHI-CUADRADO Y F DE FISHER	108
3.7	IDENTIFICACIÓN DE VALORES ATÍPICOS Y CONTROL DE CALIDAD EN MEDICIONES.....	112
4	INFERENCIA ESTADÍSTICA EN LA INVESTIGACIÓN SANITARIA	118
4.1	CONCEPTOS DE ESTIMACIÓN PUNTUAL E INTERVALOS DE CONFIANZA	118
4.2	PRUEBAS DE HIPÓTESIS: FORMULACIÓN Y TIPOS DE ERROR	122
4.3	PRUEBA T PARA MUESTRAS INDEPENDIENTES Y RELACIONADAS	127
4.4	PRUEBA CHI-CUADRADO EN ANÁLISIS DE PROPORCIONES CLÍNICAS.....	131
4.5	ANOVA EN COMPARACIÓN DE TRATAMIENTOS O GRUPOS BIOLÓGICOS.....	136
4.6	CORRELACIÓN Y REGRESIÓN LINEAL SIMPLE	139
4.7	INTERPRETACIÓN DE RESULTADOS ESTADÍSTICOS EN ESTUDIOS CLÍNICOS.....	143
5	BIOESTADÍSTICA APLICADA AL LABORATORIO CLÍNICO	148
5.1	VARIABILIDAD BIOLÓGICA Y ANALÍTICA EN PRUEBAS DE LABORATORIO	148
5.2	CONTROL DE CALIDAD INTERNO Y EXTERNO	152
5.3	INTERVALOS DE REFERENCIA Y VALORES CRÍTICOS.....	156
5.4	VALIDACIÓN ESTADÍSTICA DE MÉTODOS ANALÍTICOS	159
5.5	CONCORDANCIA DIAGNÓSTICA: SENSIBILIDAD, ESPECIFICIDAD Y VALORES PREDICTIVOS.....	163
5.6	CURVAS ROC Y ANÁLISIS DE DESEMPEÑO DE PRUEBAS DIAGNÓSTICAS	166
5.7	APLICACIONES PRÁCTICAS EN BIOQUÍMICA, HEMATOLOGÍA Y MICROBIOLOGÍA.....	170
6	ANÁLISIS EPIDEMIOLÓGICO Y ESTUDIOS CLÍNICOS	179
6.1	TIPOS DE ESTUDIOS EPIDEMIOLÓGICOS: DESCRIPTIVOS, ANALÍTICOS Y EXPERIMENTALES	179
6.2	MEDIDAS DE FRECUENCIA: INCIDENCIA Y PREVALENCIA	183
6.3	MEDIDAS DE ASOCIACIÓN: RIESGO RELATIVO Y ODDS RATIO	188
6.4	SESGOS Y FACTORES DE CONFUSIÓN EN LA INVESTIGACIÓN SANITARIA	192

6.5	PRUEBAS ESTADÍSTICAS APLICADAS A ESTUDIOS DE COHORTES Y CASOS-CONTROL..	195
6.6	META-ANÁLISIS Y SÍNTESIS DE RESULTADOS DE INVESTIGACIÓN	199
6.7	APLICACIONES PRÁCTICAS EN VIGILANCIA EPIDEMIOLÓGICA Y SALUD PÚBLICA	202
7	INTELIGENCIA ARTIFICIAL Y APRENDIZAJE AUTOMÁTICO EN BIOESTADÍSTICA	208
7.1	INTRODUCCIÓN AL APRENDIZAJE AUTOMÁTICO EN CIENCIAS DE LA SALUD	208
7.2	TIPOS DE ALGORITMOS: SUPERVISADOS, NO SUPERVISADOS Y DE REFUERZO..	212
7.3	USO DE REDES NEURONALES EN DIAGNÓSTICO MÉDICO AUTOMATIZADO	216
7.4	MODELOS PREDICTIVOS DE RIESGO CLÍNICO CON RANDOM FOREST Y SVM ..	219
7.5	MINERÍA DE DATOS CLÍNICOS Y BIG DATA EN SALUD	223
7.6	EVALUACIÓN DE MODELOS: PRECISIÓN, SENSIBILIDAD Y ESPECIFICIDAD.....	227
7.7	MODELO ALGORÍTMICO INTEGRAL APLICADO A LA BIOESTADÍSTICA Y LAS CIENCIAS DE LA SALUD	231
8	REFERENCIAS BIBLIOGRÁFICAS	249

PRÓLOGO

La formación en ciencias de la salud exige, cada vez con mayor claridad, una comprensión sólida y aplicada de los métodos cuantitativos que sustentan la investigación, la práctica clínica y la toma de decisiones sanitarias. En este escenario, la bioestadística se entiende no solo como un conjunto de técnicas matemáticas, sino como un lenguaje científico que permite interpretar la variabilidad biológica, evaluar la incertidumbre y transformar datos en información relevante para la salud individual y colectiva. El presente libro se estructura para responder a esta necesidad formativa desde una perspectiva rigurosa, didáctica y coherente con las exigencias actuales de la educación superior y del ejercicio profesional.

El contenido se organiza de manera progresiva, desde los fundamentos del razonamiento estadístico hasta aplicaciones avanzadas en investigación clínica, laboratorio, epidemiología e inteligencia artificial. El enfoque adoptado prioriza la aplicación práctica de los conceptos, estableciendo una relación directa entre los procedimientos estadísticos y situaciones reales del ámbito sanitario, lo que favorece una comprensión funcional y contextualizada. Esta orientación resulta especialmente pertinente para estudiantes y profesionales de disciplinas como medicina, enfermería, laboratorio clínico, salud pública y áreas afines, quienes requieren herramientas analíticas aplicables con criterio técnico y responsabilidad profesional.

Un elemento distintivo del texto es la articulación equilibrada entre los métodos bioestadísticos tradicionales y los enfoques contemporáneos de análisis de datos. Los principios de la estadística descriptiva, inferencial y epidemiológica se presentan como la base sobre la cual se incorporan modelos predictivos y técnicas de aprendizaje automático, entendidos como extensiones metodológicas del análisis cuantitativo en salud. Esta integración permite comprender la continuidad conceptual entre distintos enfoques analíticos y evita una visión fragmentada del tratamiento de la información sanitaria.

Asimismo, se enfatiza la interpretación crítica de los resultados estadísticos, reconociendo que el valor del análisis no se limita al cálculo numérico, sino que depende de su significado clínico y epidemiológico. La evaluación de la incertidumbre, la identificación de limitaciones metodológicas y la comunicación clara de la evidencia constituyen ejes transversales del texto, alineados con los principios de la medicina basada en evidencia y la investigación científica responsable.

Este libro se concibe como un recurso académico destinado a fortalecer la formación en bioestadística aplicada a las ciencias de la salud y a apoyar el desarrollo de competencias analíticas en contextos de docencia, investigación y práctica profesional. Su finalidad es contribuir al uso reflexivo y fundamentado de la estadística como herramienta para mejorar la calidad del conocimiento científico y orientar decisiones sanitarias sustentadas en datos.

Los Autores

Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

CAPÍTULO I

Fundamentos de la Bioestadística en las Ciencias de la Salud

AUTOR: Castro Jalca Jazmín Elena



1 Fundamentos de la Bioestadística en las Ciencias de la Salud

1.1 Concepto y evolución histórica de la bioestadística

La bioestadística es el campo que aplica métodos estadísticos al estudio de fenómenos biológicos, biomédicos y sanitarios, con el fin de describir patrones, cuantificar incertidumbre y sustentar inferencias sobre poblaciones a partir de datos observados. En ciencias de la salud, su alcance incluye desde la descripción de parámetros fisiológicos y resultados de laboratorio, hasta la evaluación de tratamientos, la vigilancia epidemiológica y la estimación de riesgo. Esta definición operativa coincide con la tradición académica que presenta la bioestadística como un conjunto de herramientas para diseñar estudios, analizar datos y comunicar evidencia cuantitativa con validez científica.

En términos didácticos, resulta útil entender la bioestadística como una disciplina con doble naturaleza:

- a) Metodológica, porque provee procedimientos de muestreo, estimación, contraste y modelado.
- b) Aplicada, porque sus técnicas se ajustan a necesidades reales de la práctica clínica, el laboratorio clínico y la salud pública.

Esta orientación aplicada se fortaleció a medida que la medicina y la epidemiología incorporaron diseños comparativos, registro sistemático de eventos y evaluación cuantitativa de intervenciones (Sackett et al., 1996).

1.1.1 Nacimiento del Razonamiento cuantitativo en salud

Antes de la formalización estadística, diversas sociedades mantuvieron registros de nacimientos y defunciones con fines administrativos. Sin embargo, el paso hacia el análisis cuantitativo asociado a salud se asocia con el estudio de las listas de mortalidad en Inglaterra. En 1662, John Graunt publicó *Natural and Political Observations Made upon the Bills of Mortality*, obra reconocida como un antecedente mayor de la epidemiología cuantitativa por su análisis sistemático de datos de mortalidad, patrones por edad y tendencias temporales. Investigaciones históricas en epidemiología destacan ese texto como un hito temprano en el uso de datos poblacionales para producir inferencias sobre salud y supervivencia (Morabia, 2013).

Una forma simple de ilustrar ese avance, compatible con la enseñanza inicial, es introducir el concepto de tasa, que expresa eventos por unidad de población y tiempo. Por ejemplo, una tasa de mortalidad general puede formularse como:

$$\text{Tasa de mortalidad} = \frac{\text{Nº de defunciones en un periodo}}{\text{Población en el mismo periodo}} \times k$$

donde *k* suele ser 1,000 o 100,000, según el estándar de reporte. Aunque este formalismo se consolidó posteriormente, su lógica está alineada con el tipo de cuantificación que impulsó el análisis de mortalidad desde el siglo XVII.

1.1.2 Estadística vital y salud pública

En el siglo XIX, el crecimiento urbano e industrial incrementó el interés por medir mortalidad, morbilidad y condiciones sanitarias. En ese periodo, Farr (2014) realizó aportes fundamentales a la historia de la salud pública mediante su labor en los registros civiles, la elaboración de tablas de vida y la clasificación de causas de muerte. Fuentes históricas destacan su trabajo en el General Register Office del Reino Unido y su influencia en el desarrollo de la estadística médica.

La institucionalización de registros civiles y censos más completos permitió pasar de conteos locales a comparaciones territoriales y temporales. En términos pedagógicos, este punto marca una transición: el dato deja de ser una anotación aislada y se convierte en una serie susceptible de comparación, control y seguimiento. Un ejemplo que se puede incorporar como ejercicio es el cálculo de esperanza de vida con tablas de vida, que se apoya en probabilidades de muerte por intervalos de edad y en la suma de años-persona vividos por una cohorte hipotética. Textos históricos y compilaciones de la obra de Farr muestran la relevancia de la estadística vital como fundamento para indicadores sanitarios.

Tabla 1-1. Hitos históricos y aportes formativos para la bioestadística

Periodo	Referente	Aporte cuantitativo	Relevancia formativa en ciencias de la salud
1662	John Graunt	Análisis de datos de mortalidad y patrones poblacionales	Introduce lectura de series de mortalidad y razonamiento con tasas

1839– 1879 (siglo XIX)	William Farr	Estadística vital, tablas de vida, clasificación de causas de muerte	Fortalece indicadores sanitarios comparables y vigilancia poblacional
1850s– 1860s	Florence Nightingale	Visualización estadística de mortalidad y reforma sanitaria	Une evidencia numérica con comunicación visual para decisiones sanitarias
1935	R. A. Fisher	Diseño experimental, aleatorización, ANOVA, prueba de significación	Fundamento de ensayos y comparación entre grupos
1948	MRC Streptomycin Trial	Ensayo aleatorizado moderno y evaluación terapéutica	Base para investigación clínica controlada
1996	Sackett y cols.	Medicina basada en evidencia y uso explícito de investigación	Integra evidencia cuantitativa en decisiones clínicas

1.1.3 Visualización y reforma sanitaria: Florence Nightingale y el uso persuasivo de datos

La contribución de Florence Nightingale se destaca por el uso de gráficos para mostrar que gran parte de las muertes en la Guerra de Crimea se asociaban a enfermedades prevenibles y condiciones sanitarias deficientes. Sus diagramas circulares (“polar area diagrams”) son citados como ejemplos tempranos de comunicación estadística orientada a reforma de sistemas hospitalarios. La documentación histórica y la disponibilidad pública de la figura original respaldan esta relación entre visualización y decisiones sanitarias.

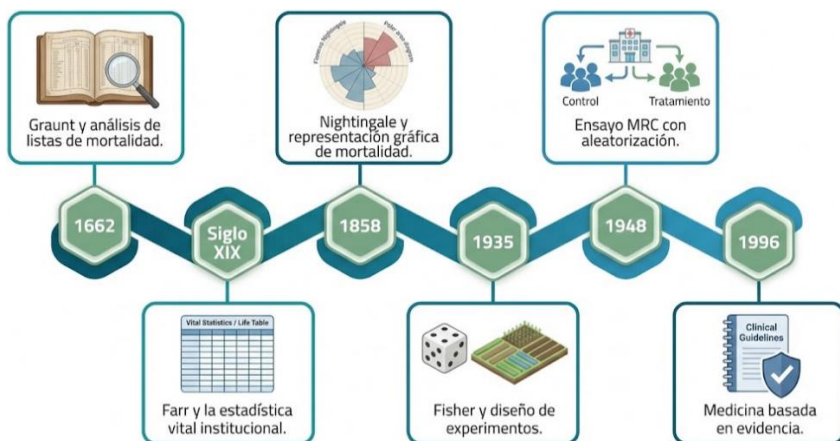


Figura 1-1. Línea temporal de la bioestadística aplicada a salud (1662-1996).

1.1.4 Estadística moderna y diseño experimental

El siglo XX consolidó la estadística matemática aplicada a investigación científica. La obra *The Design of Experiments* de Fisher (1971) sistematizó principios de aleatorización, replicación y control de variabilidad, esenciales para comparar tratamientos y reducir inferencias sesgadas. Fuentes académicas y reproducciones del texto señalan su influencia en el desarrollo del análisis de varianza y del enfoque de prueba de significación, con impacto directo en investigación biomédica.

Un aspecto útil para docencia es mostrar, de forma muy simple, qué significa comparar medias entre grupos con variabilidad. Sin entrar todavía al capítulo de inferencia, puede colocarse una fórmula descriptiva como la media muestral:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

y la varianza muestral:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Estas expresiones apoyan la idea de que la evidencia numérica requiere cuantificar dispersión, no solo valores centrales, lo cual se vuelve decisivo cuando se comparan grupos clínicos o resultados de laboratorio.

El desarrollo del contraste de hipótesis también se enriqueció con aportes posteriores sobre criterios de decisión y clasificación de errores. La discusión sobre errores tipo I y tipo II y la evolución de los enfoques de prueba se trata en metodología contemporánea orientada a enseñanza, que explica diferencias entre tradiciones fisherianas y neyman-pearsonianas en el aprendizaje de procedimientos de prueba (Crofton, 2006).

1.1.5 Ensayos clínicos y estandarización: el legado del ensayo MRC de estreptomicina

En investigación clínica, un punto de referencia histórico es el ensayo aleatorizado del Medical Research Council (MRC) sobre estreptomicina para tuberculosis pulmonar (1948). Revisiones históricas destacan su valor por el uso de asignación aleatoria y procedimientos controlados, que abrieron el camino a los ensayos clínicos contemporáneos (Crofton, 2006).

En docencia, este hito permite explicar por qué la aleatorización equilibra factores pronósticos conocidos y desconocidos entre grupos, y reduce distorsiones sistemáticas. Además, facilita introducir el lenguaje de comparación basado en riesgos, proporciones de respuesta o cambios promedio, sin necesidad de desarrollar aún pruebas formales.

1.1.6 Medicina basada en evidencia y expansión de la bioestadística en la formación sanitaria

En 1996, Sackett y colaboradores definieron la medicina basada en evidencia como el uso consciente y explícito de la mejor evidencia disponible, junto con la experiencia clínica, para tomar decisiones sobre la atención al paciente. Ese trabajo se convirtió en una referencia formativa en educación médica, al reforzar la necesidad de competencias estadísticas para interpretar ensayos, revisiones sistemáticas y guías clínicas (Altman & Bland, 2011).

Este cambio incrementó la demanda educativa: el personal de salud debía comprender medidas de efecto, intervalos de confianza, significación

estadística y sesgos, además de adquirir habilidades para la lectura crítica de artículos. En ese marco, recursos metodológicos de divulgación científica también contribuyeron a mejorar la enseñanza de la interpretación estadística en medicina, con notas orientadas al lector clínico y al investigador aplicado.

1.1.7 Bioestadística en el siglo XXI

En décadas recientes, la bioestadística se ha ampliado hacia biología molecular, genómica, bioinformática y análisis de grandes bases de datos clínicos. Esta disciplina incorporó métodos de modelado predictivo y estrategias computacionales para gestionar grandes volúmenes de información sanitaria.

En la enseñanza universitaria, este crecimiento plantea un reto: formar profesionales capaces de comprender supuestos, interpretar resultados y reconocer limitaciones de datos clínicos reales, como calidad de registro, medición, pérdidas y sesgo de selección. Esta formación no se reduce a memorizar pruebas; exige construir criterio para seleccionar métodos y comunicar hallazgos de forma responsable.

Tabla 1-2. Evolución del enfoque bioestadístico según el tipo de pregunta sanitaria

Tipo de pregunta	Ejemplo en ciencias de la salud	Método asociado	Resultado típico
Descriptiva	¿Cuál es la distribución de glucosa en ayuno en adultos?	Medidas descriptivas, percentiles	Media/mediana, DE, rangos
Comparativa	¿Un fármaco reduce presión arterial frente a control?	Diseño experimental y comparación	Diferencia de medias o proporciones
Asociativa	¿Tabaquismo se asocia con EPOC?	Medidas de asociación y modelado	RR/OR, regresión

Diagnóstica	¿Qué tan bien clasifica una prueba?	Sensibilidad/especificidad, ROC	AUC, puntos de corte
Decisional	¿Qué evidencia sustenta una guía clínica?	Revisiones, lectura crítica	Recomendaciones basadas en evidencia

Nota. La tabla muestra cómo se relaciona preguntas sanitarias frecuentes con métodos cuantitativos y resultados típicos usados en práctica clínica y salud pública.

1.2 Importancia de la bioestadística en la investigación biomédica y clínica

La investigación biomédica y clínica se apoya en datos que provienen de mediciones fisiológicas, resultados de laboratorio, imágenes, cuestionarios, historias clínicas y sistemas de vigilancia. En todos esos escenarios existe variabilidad: dos pacientes con el mismo diagnóstico pueden responder de forma distinta a un tratamiento, un analizador clínico puede mostrar dispersión aun con controles estables, y una muestra puede no representar perfectamente a la población. La bioestadística aporta métodos para manejar esa variabilidad, estimar magnitudes con incertidumbre cuantificada y sostener decisiones con criterios explícitos de validez. Esta orientación se integra con el enfoque de medicina basada en evidencia, que propone combinar la mejor evidencia de investigación con la experiencia clínica y las preferencias del paciente (Sackett et al., 1996).

1.2.1 Razones científicas y clínicas que justifican su uso

1. Diseño correcto de estudios y prevención de conclusiones erróneas

Un diseño deficiente genera resultados que pueden ser estadísticamente llamativos, pero clínicamente engañosos. La bioestadística orienta la elección del diseño, el muestreo, el tamaño muestral, la definición de variables, el plan de análisis y el manejo de datos faltantes. Las guías internacionales para ensayos clínicos establecen que la metodología estadística debe estar planificada desde

el protocolo, incluyendo objetivos, variables primarias, criterios de análisis y manejo de desviaciones del protocolo (ICH E9) (EMA, 1998).

2. Cuantificación de la incertidumbre

En salud no basta con reportar un valor promedio o una proporción. Es necesario acompañarlos de una medida de precisión que permita valorar la confiabilidad de la estimación. En investigación clínica, los intervalos de confianza se recomiendan como parte del reporte por su capacidad para expresar magnitud y precisión, complementando el valor p. Comunican tanto la estimación como su incertidumbre para evitar interpretaciones simplistas.

3. Interpretación crítica y transparencia del reporte

La calidad de la evidencia depende también de cómo se reporta. En ensayos clínicos aleatorizados, las guías CONSORT establecen elementos mínimos para reportar métodos y resultados, incluyendo asignación, flujo de participantes y análisis, para que el lector pueda evaluar sesgos y aplicabilidad (Hopewell et al., 2025). En estudios observacionales, STROBE promueve reportes completos de métodos estadísticos, control de confusión, manejo de datos faltantes y justificación del tamaño del estudio, facilitando la evaluación crítica (Elm et al., 2007).

4. Toma de decisiones clínicas, regulatorias y sanitarias

Los ensayos y estudios observacionales generan evidencia que impacta guías clínicas, políticas de salud y aprobación de tecnologías sanitarias. Por ello, documentos regulatorios y de buenas prácticas exigen coherencia entre pregunta clínica, diseño, análisis y reporte. ICH E9 enfatiza que el análisis estadístico es parte integral de la evaluación de eficacia y seguridad.

1.2.2 Bioestadística y calidad metodológica

En la enseñanza, es útil vincular la bioestadística con tres criterios evaluativos:

- **Validez interna:** grado en que el resultado se atribuye a la intervención o exposición, no a sesgos o confusión. CONSORT promueve reportar elementos que protegen la validez interna como aleatorización y cegamiento cuando corresponda.
- **Validez externa:** posibilidad de aplicar resultados a poblaciones similares. STROBE impulsa describir participantes, fuentes de datos y criterios de inclusión para valorar generalización.
- **Precisión:** estrechez del intervalo de confianza y estabilidad de la estimación. ICH E9 desarrolla principios para estimación y evaluación estadística de resultados en ensayos.

Tabla 1-3. Aportes directos de la bioestadística según fase de investigación clínica

Fase	Actividades	Aporte bioestadístico	Producto esperado
Planteamiento	Objetivos, hipótesis, variables	Definición operacional y plan analítico	Protocolo coherente
Diseño	Selección de diseño, aleatorización, control	Minimiza sesgos y confusión	Validez interna mejorada
Recolección	Instrumentos, calidad de datos	Control de medición y consistencia	Base de datos depurada
Análisis	Estimación, pruebas, modelos	Cuantifica efecto e incertidumbre	Estimaciones e IC
Reporte	Artículo y anexos	Transparencia y reproducibilidad	Reporte conforme a guías

Nota. La tabla relaciona fases del estudio con decisiones estadísticas que elevan validez, precisión y calidad del reporte.

La Figura 1.2 ilustra el ciclo de la evidencia cuantitativa en las ciencias de la salud, mostrando de manera integrada las etapas que vinculan la investigación biomédica con la toma de decisiones clínicas y sanitarias. El proceso se inicia

con la formulación de una pregunta clínica o sanitaria claramente definida, la cual orienta el diseño metodológico del estudio. Posteriormente, la recolección de datos y su control de calidad garantizan la confiabilidad de la información obtenida, permitiendo un análisis estadístico previamente planificado y acorde con los objetivos del estudio. Los resultados derivados de dicho análisis deben ser interpretados desde una perspectiva clínica, considerando explícitamente la incertidumbre asociada a las estimaciones.



Figura 1-2. Ciclo de evidencia cuantitativa en ciencias de la salud.

1.2.3 Fórmulas esenciales que conectan evidencia con decisión

A nivel introductorio, conviene introducir fórmulas que aparecen repetidamente en artículos biomédicos.

a) Intervalo de confianza para una media (cuando n es moderado y la variable es aproximadamente normal)

$$IC_{95\%} = \bar{x} \pm t_{0.975, n-1} \cdot \frac{s}{\sqrt{n}}$$

Donde $\bar{x}$ es la media, s la desviación estándar, n el tamaño muestral, y t el valor crítico de la distribución t .

b) Intervalo de confianza para una proporción

$$IC_{95\%} = \hat{p} \pm 1.96 \cdot \sqrt{\frac{\hat{p}(1 - \hat{p})}{n}}$$

Estas expresiones permiten enseñar que toda estimación debe acompañarse de precisión, idea coherente con la práctica recomendada para reportes de investigación clínica y sanitaria.

1.2.3.1 Ejemplo de la aplicación biomédica y clínica

Ejercicio 1. Estimación e interpretación clínica con intervalo de confianza (media de glucosa)

Planteamiento:

En un estudio transversal se mide glucosa en ayunas (mg/dL) en $n = 25$ adultos. Se obtuvo $\bar{x} = 102$ y $s = 12$. Calcule el intervalo de confianza del 95% e interprételo.

Solución:

1. Fórmula:

$$IC_{95\%} = \bar{x} \pm t_{0.975, n-1} \cdot \frac{s}{\sqrt{n}}$$

2. Grados de libertad: $n - 1 = 24$. Valor aproximado: $t_{0.975,24} = 2.064$.

3. Error estándar:

$$\frac{s}{\sqrt{n}} = \frac{12}{\sqrt{25}} = \frac{12}{5} = 2.4$$

4. Margen:

$$2.064 \times 2.4 = 4.9536 = 4.95$$

5. Intervalo:

$$IC_{95\%} = 102 \pm 4.95 \Rightarrow (97.05, 106.95)$$

Interpretación:

La media poblacional de glucosa en ayunas se estima entre 97.1 y 107.0 mg/dL con 95% de confianza. El intervalo expresa precisión; si el objetivo clínico fuera comparar con un valor de referencia, esta amplitud es informativa para decidir si se requiere mayor tamaño muestral.

Ejercicio 2. Comparación de proporciones y lectura de evidencia en un ensayo

Planteamiento:

En un ensayo clínico, 120 pacientes reciben Tratamiento A y 120 reciben control. Se observa mejoría clínica en 78 del grupo A y 60 del control.

1. Estime la diferencia de proporciones.
2. Calcule un intervalo aproximado del 95% para la diferencia.
3. Interprete el resultado.

Solución:

1. Proporciones:

$$\hat{p}_A = \frac{78}{120} = 0.65, \hat{p}_C = \frac{60}{120} = 0.50$$

Diferencia:

$$\Delta = \hat{p}_A - \hat{p}_C = 0.65 - 0.50 = 0.15$$

2. Error estándar aproximado para diferencia de proporciones:

$$SE = \sqrt{\frac{\hat{p}_A(1 - \hat{p}_A)}{n_A} + \frac{\hat{p}_C(1 - \hat{p}_C)}{n_C}}$$

$$SE = \sqrt{\frac{0.65(0.35)}{120} + \frac{0.50(0.50)}{120}} = \sqrt{\frac{0.2275}{120} + \frac{0.25}{120}}$$

$$= \sqrt{0.0018958 + 0.0020833} = \sqrt{0.0039791} = 0.0631$$

3. Intervalo aproximado 95%:

$$IC_{95\%} \approx \Delta \pm 1.96 \times SE = 0.15 \pm 1.96 \times 0.0631 = 0.15 \pm 0.1237$$

$$\Rightarrow (0.0263, 0.2737)$$

Interpretación:

El tratamiento A incrementa la proporción de mejoría en 15 puntos porcentuales; el intervalo sugiere que el aumento real podría estar entre 2.6 y 27.4 puntos. Esta forma de presentación se alinea con reportes orientados a magnitud y precisión, coherentes con buenas prácticas metodológicas y de reporte de ensayos.

Ejercicio 3. Tamaño de muestra básico para estimar una proporción (vigilancia sanitaria)

Planteamiento:

Se desea estimar la prevalencia de hipertensión en una comunidad con error máximo $d = 0.05$ y 95% de confianza. Si no se conoce la prevalencia, use $p = 0.50$. Calcule el tamaño de muestra mínimo.

Solución:

Fórmula clásica:

$$n = \frac{Z_{0.975}^2 p(1-p)}{d^2}$$

Con $Z_{0.975} = 1.96$, $p = 0.50$, $d = 0.05$:

$$n = \frac{1.96^2 \times 0.5(0.5)}{0.05^2} = \frac{3.8416 \times 0.25}{0.0025} = \frac{0.9604}{0.0025} = 384.16$$

Se redondea hacia arriba: $n = 385$.

Interpretación:

Con 385 personas se logra estimar la prevalencia con error máximo 5% y 95% de confianza. Esta planificación responde a principios metodológicos de investigación sanitaria que exigen justificar el tamaño de estudio y su precisión esperada.

Tabla 1-4. Errores frecuentes al interpretar resultados estadísticos en clínica

Situación	Error frecuente	Consecuencia	Enfoque recomendado
Valor p significativo	Confundir significación con relevancia clínica	Decisiones exageradas	Reportar efecto e IC
Muestra pequeña	Concluir “no efecto” por p mayor a 0.05	Falso negativo	Valorar potencia y precisión
Observacional	Interpretar asociación como causalidad	Inferencias inválidas	Control de confusión y reporte STROBE
Ensayo	Reporte incompleto del flujo	Sesgo por pérdidas	Seguir CONSORT

Nota. Errores interpretativos habituales que afectan decisiones; se proponen prácticas alineadas con guías metodológicas y de reporte.

La bioestadística sostiene la investigación biomédica y clínica porque ordena el proceso científico: transforma preguntas sanitarias en diseños evaluables, convierte datos en estimaciones con incertidumbre cuantificada y exige comunicación transparente para evaluación crítica. Este enfoque es consistente

con la medicina basada en evidencia y con guías internacionales que orientan el reporte y la evaluación de ensayos y estudios observacionales.

1.3 Tipos de variables: cualitativas, cuantitativas y sus escalas de medición

El análisis estadístico en las ciencias de la salud se fundamenta, de manera estricta, en la correcta identificación y clasificación de las variables que intervienen en un estudio. Una variable puede definirse como cualquier característica observable o medible en un individuo, muestra o población, cuyo valor puede variar entre sujetos o a lo largo del tiempo. En investigación biomédica y clínica, las variables representan fenómenos biológicos, clínicos, conductuales, ambientales o sociales, y su adecuada definición condiciona el diseño del estudio, el tipo de análisis estadístico y la interpretación de los resultados.

Desde el punto de vista metodológico, la clasificación de las variables no constituye un ejercicio meramente descriptivo. Al contrario, determina qué procedimientos estadísticos son válidos, qué resúmenes numéricos pueden calcularse y qué tipo de inferencias pueden realizarse sin introducir errores conceptuales. Diversos manuales de bioestadística aplicada a la salud coinciden en que una proporción significativa de errores en investigaciones clínicas se origina en la incorrecta identificación del tipo de variable y de su escala de medición (Altman & Bland, 2011; Motulsky H., 2018).

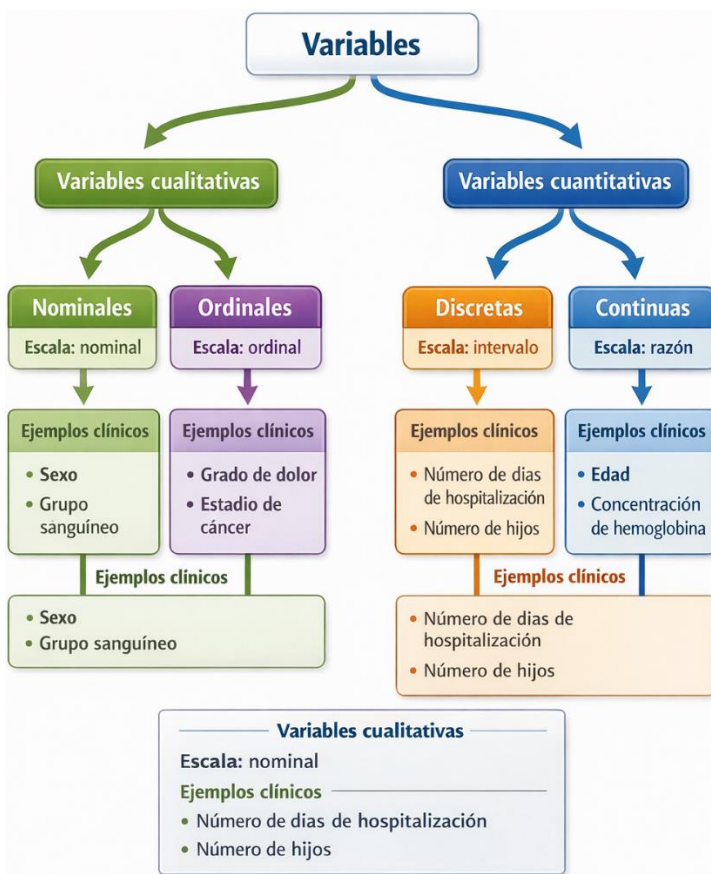


Figura 1-3. Clasificación de variables en ciencias de la salud. Esta figura facilita la comprensión visual de la clasificación y su aplicación práctica en estudios biomédicos y clínicos

1.3.1 Variables cualitativas o categóricas

Las variables cualitativas, también denominadas categóricas, describen atributos, cualidades o categorías que no se expresan mediante valores numéricos con significado aritmético. En el ámbito sanitario, estas variables son frecuentes y esenciales, ya que muchas características clínicas no pueden cuantificarse de forma continua, sino que se clasifican en categorías discretas (W. W. Daniel & Chad L, 2018).

Ejemplos comunes de variables cualitativas en ciencias de la salud incluyen el sexo biológico, el grupo sanguíneo, el estado vital, el tipo de diagnóstico, la presencia o ausencia de una enfermedad, y la clasificación de resultados de pruebas diagnósticas como positivo o negativo.

Las variables cualitativas se subdividen en dos grandes grupos según la existencia o no de un orden lógico entre sus categorías.

a) Variables cualitativas nominales

Las variables nominales clasifican a los individuos en categorías mutuamente excluyentes, sin un orden jerárquico intrínseco. Las categorías funcionan únicamente como etiquetas. En estas variables no es posible establecer relaciones de mayor o menor entre los valores observados (Pagano et al., 2022).

Ejemplos en salud incluyen:

- Grupo sanguíneo: A, B, AB, O
- Tipo de patógeno: bacteriano, viral, fúngico
- Diagnóstico principal: diabetes, hipertensión, asma

En el análisis estadístico, estas variables se resumen mediante frecuencias absolutas, relativas y porcentajes, y se analizan mediante pruebas basadas en conteos, como la prueba chi cuadrado o pruebas exactas, según corresponda (Agresti, 2019).

b) Variables cualitativas ordinales

Las variables ordinales presentan categorías con un orden lógico, aunque la distancia entre las categorías no es cuantificable de forma exacta. Este tipo de variables es muy común en escalas clínicas, cuestionarios y clasificaciones pronósticas.

Ejemplos frecuentes son:

- Grado de dolor: leve, moderado, severo
- Estadio clínico de una enfermedad: I, II, III, IV

- Nivel de dependencia funcional: independiente, parcialmente dependiente, dependiente

Aunque estas variables pueden codificarse con números, dichos valores no deben interpretarse como cantidades aritméticas. Por esta razón, las medidas de tendencia central recomendadas suelen ser la mediana y los percentiles, y los análisis inferenciales suelen apoyarse en pruebas no paramétricas (Conover, 1999).

Tabla 1-5. Variables cualitativas en ciencias de la salud

Tipo	Característica	Ejemplo clínico	Resumen estadístico habitual
Nominal	Sin orden entre categorías	Grupo sanguíneo	Frecuencias y porcentajes
Nominal dicotómica	Dos categorías	Presencia de enfermedad	Proporciones
Ordinal	Categorías ordenadas	Grado de dolor	Mediana y percentiles

Nota. La tabla muestra la clasificación de variables categóricas según presencia de orden y formas habituales de descripción en estudios sanitarios.

1.3.2 Variables cuantitativas o numéricas

Las variables cuantitativas representan características que pueden medirse numéricamente y en las cuales las operaciones aritméticas tienen sentido. En investigación biomédica, estas variables permiten describir magnitudes fisiológicas, bioquímicas y clínicas con alto nivel de precisión (Rosner Bernard, 2016).

Ejemplos comunes incluyen edad, peso corporal, presión arterial, concentración de glucosa, frecuencia cardíaca y niveles hormonales.

Las variables cuantitativas se clasifican según la forma en que toman valores en dos categorías principales.

a) Variables cuantitativas discretas

Las variables discretas solo pueden asumir valores enteros, resultado de un conteo. No admiten valores intermedios entre dos números consecutivos.

Ejemplos:

- Número de consultas médicas en un año
- Número de episodios de infección
- Número de hospitalizaciones

Estas variables suelen resumirse mediante medias, medianas o tasas, y se analizan con modelos específicos cuando presentan distribución asimétrica, como modelos de Poisson o binomiales negativos en estudios epidemiológicos.

b) Variables cuantitativas continuas

Las variables continuas pueden tomar cualquier valor real dentro de un intervalo determinado, dependiendo de la precisión del instrumento de medición. Son predominantes en el laboratorio clínico y en estudios fisiológicos.

Ejemplos:

- Concentración sérica de colesterol
- Presión arterial sistólica
- Índice de masa corporal

Estas variables permiten el uso de una amplia gama de técnicas estadísticas, siempre que se verifiquen los supuestos de los modelos aplicados.

Tabla 1-6. Variables cuantitativas y ejemplos biomédicos

Tipo	Definición	Ejemplo	Observación metodológica
Discreta	Valores enteros contables	Número de eventos	Puede requerir modelos específicos

Continua	Valores reales medibles	Glucosa en sangre	Depende de precisión instrumental
----------	-------------------------	-------------------	-----------------------------------

Nota. La tabla muestra la diferenciación entre variables numéricas según forma de medición y consecuencias analíticas en investigación sanitaria.

1.3.3 Escalas de medición y su relevancia estadística

Además de distinguir entre variables cualitativas y cuantitativas, es imprescindible reconocer la escala de medición, ya que esta define las operaciones matemáticas y los análisis estadísticos que pueden realizarse de manera válida. La clasificación clásica reconoce cuatro escalas de medición: nominal, ordinal, de intervalo y de razón (Stevens, 1946).

- a. Escala nominal**
Corresponde a variables cualitativas sin orden. Permite únicamente clasificar y contar.
- b. Escala ordinal**
Corresponde a variables con orden, pero sin distancias iguales entre categorías.
- c. Escala de intervalo**
Se aplica a variables cuantitativas donde las diferencias entre valores son significativas, pero el cero es arbitrario. Un ejemplo clásico es la temperatura corporal en grados Celsius.
- d. Escala de razón**
Presenta intervalos iguales y un cero absoluto, lo que permite todas las operaciones aritméticas. Ejemplos incluyen peso, talla y concentraciones bioquímicas.

Tabla 1-7. Escalas de medición y ejemplos en salud

Escala	Característica	Ejemplo	Operaciones válidas
Nominal	Clasificación sin orden	Grupo sanguíneo	Conteo
Ordinal	Orden sin distancia fija	Estadio clínico	Comparación

Intervalo	Distancias iguales, cero arbitrario	Temperatura	Suma y resta
Razón	Cero absoluto	Peso corporal	Todas

Nota. La tabla muestra la relación entre escala de medición y tipo de operaciones estadísticas permitidas en investigación biomédica.

1.3.4 Implicaciones estadísticas y errores frecuentes

La confusión entre tipos de variables y escalas de medición conduce a errores metodológicos importantes. Un ejemplo habitual es calcular promedios en variables ordinales o aplicar pruebas paramétricas sin verificar la naturaleza de los datos. Manuales de estadística médica advierten que tales prácticas pueden producir conclusiones inválidas, aun cuando los cálculos sean correctos desde el punto de vista matemático.

Por ello, la enseñanza de la bioestadística en ciencias de la salud debe insistir en la identificación previa del tipo de variable antes de seleccionar cualquier procedimiento analítico.

Ejercicio aplicado 1. Clasificación de variables clínicas

Planteamiento:

En un estudio hospitalario se recolectan las siguientes variables: edad, sexo, concentración de hemoglobina, grado de dolor postoperatorio, número de días de hospitalización.

Solución:

- Edad: cuantitativa continua, escala de razón
- Sexo: cualitativa nominal
- Hemoglobina: cuantitativa continua, escala de razón
- Grado de dolor: cualitativa ordinal
- Días de hospitalización: cuantitativa discreta, escala de razón

Interpretación:

La clasificación adecuada permite seleccionar medidas descriptivas y pruebas estadísticas consistentes con la naturaleza de cada variable.

Ejercicio aplicado 2. Selección del resumen estadístico adecuado**Planteamiento:**

Indique el resumen estadístico más apropiado para las siguientes variables:

- a) Grupo sanguíneo
- b) Presión arterial sistólica
- c) Estadio de cáncer

Solución:

- a) Frecuencias y porcentajes
- b) Media y desviación estándar
- c) Mediana y rango intercuartílico

Propuesta de figura esquemática en forma de diagrama jerárquico que muestre:

- Variables cualitativas: nominales y ordinales
- Variables cuantitativas: discretas y continuas
- Correspondencia con escalas de medición y ejemplos clínicos

La correcta identificación de los tipos de variables y de sus escalas de medición constituye un requisito indispensable para la validez del análisis estadístico en ciencias de la salud. Esta etapa inicial orienta el diseño del estudio, el resumen de los datos y la selección de métodos inferenciales adecuados, evitando errores conceptuales que comprometan la interpretación clínica y científica de los resultados (Pagano et al., 2022; Rosner Bernard, 2016).

1.4 Población, muestra y parámetros estadísticos

La investigación en ciencias de la salud se orienta a generar conocimiento válido sobre fenómenos que afectan a grupos humanos, comunidades o sistemas biológicos. Sin embargo, en la práctica clínica y sanitaria rara vez es posible estudiar a todos los individuos que conforman el conjunto de interés. Por esta razón, la bioestadística introduce los conceptos de población y muestra como fundamentos metodológicos que permiten obtener conclusiones

confiables a partir de observaciones parciales, siempre que se apliquen procedimientos adecuados de selección y análisis.

El uso correcto de estos conceptos resulta decisivo para la validez de los estudios biomédicos. Una definición imprecisa de la población o una muestra mal seleccionada comprometen la interpretación clínica de los resultados y limitan su aplicación en la práctica sanitaria. Diversos textos de estadística médica coinciden en que los errores en esta etapa inicial son una de las principales fuentes de sesgo en la investigación clínica y epidemiológica (Altman, 1991; Pagano y Gauvreau, 2018).



Figura 1-4. Relación entre población, muestra y estadísticos. Esta representación facilita la comprensión visual del proceso inferencial en bioestadística aplicada a la salud.

1.4.1 Población estadística en ciencias de la salud

La población estadística se define como el conjunto total de individuos, unidades o elementos que cumplen con características previamente establecidas y sobre los cuales se desea realizar inferencias. En el ámbito de la salud, la población puede estar conformada por personas, pacientes, muestras biológicas, registros clínicos o eventos sanitarios, dependiendo del objetivo del estudio.

Es fundamental distinguir entre población objetivo y población accesible. La población objetivo corresponde al conjunto teórico al que se desea generalizar los resultados, por ejemplo, “adultos con diabetes tipo 2 atendidos en atención primaria”. La población accesible, en cambio, está constituida por aquellos individuos que pueden ser efectivamente incluidos en el estudio, como “pacientes con diabetes tipo 2 atendidos en un hospital específico durante un periodo determinado” (W. W. Daniel & Chad L, 2018).

Esta diferenciación es especialmente relevante en estudios clínicos y epidemiológicos, donde las limitaciones logísticas, éticas y temporales condicionan el acceso a los sujetos de investigación.

1.4.2 Muestra y su función en la inferencia estadística

La muestra es un subconjunto de la población que se selecciona con el propósito de obtener información representativa del conjunto total. El principio central que justifica el uso de muestras es que, bajo ciertas condiciones, las características observadas en una muestra bien seleccionada permiten estimar con precisión los parámetros poblacionales.

En ciencias de la salud, el muestreo cumple una función doble. Por un lado, hace viable la investigación al reducir costos, tiempo y exposición de los participantes. Por otro, permite aplicar técnicas estadísticas que cuantifican la incertidumbre asociada a las estimaciones, elemento indispensable para la toma de decisiones clínicas y sanitarias.

La representatividad de la muestra no depende únicamente de su tamaño, sino también del método de selección. Muestras grandes pero sesgadas pueden

producir estimaciones más erróneas que muestras más pequeñas seleccionadas de manera adecuada.

Tabla 1-8. Diferencias conceptuales entre población y muestra en investigación sanitaria

Concepto	Definición	Ejemplo en salud	Implicación metodológica
Población	Conjunto total de interés	Pacientes con hipertensión	Marco de inferencia
Muestra	Subconjunto de la población	Pacientes evaluados	Base del análisis

Nota. La tabla muestra la diferencia entre el conjunto total de interés y el subconjunto observado para realizar inferencias estadísticas en salud.

1.4.3 Tipos de muestreo en investigación biomédica

El muestreo se refiere al procedimiento mediante el cual se seleccionan los elementos que integran la muestra. En bioestadística aplicada a la salud se distinguen, de manera general, métodos probabilísticos y no probabilísticos, cada uno con implicaciones distintas para la validez de los resultados.

a) Muestreo probabilístico

En el muestreo probabilístico, todos los individuos de la población tienen una probabilidad conocida y distinta de cero de ser seleccionados. Este enfoque permite aplicar la teoría de la probabilidad para estimar errores muestrales y construir intervalos de confianza válidos.

Entre las modalidades más utilizadas se encuentran:

- Muestreo aleatorio simple
- Muestreo sistemático
- Muestreo estratificado
- Muestreo por conglomerados

El muestreo estratificado es especialmente útil en estudios de salud cuando se desea garantizar la representación de subgrupos relevantes, como sexo, grupos etarios o niveles de severidad clínica.

b) Muestreo no probabilístico

En este tipo de muestreo, la selección de los sujetos no se basa en probabilidades conocidas. Aunque es frecuente en estudios exploratorios y en entornos clínicos con limitaciones prácticas, su principal desventaja es la dificultad para generalizar los resultados a la población objetivo.

Ejemplos comunes incluyen muestreo por conveniencia y muestreo intencional. Estos métodos requieren una interpretación prudente de los resultados y una descripción transparente de los criterios de inclusión.

Tabla 1-9. Métodos de muestreo y aplicaciones en ciencias de la salud

Tipo	Método	Uso frecuente	Limitación principal
Probabilístico	Aleatorio simple	Estudios poblacionales	Requiere marco muestral
Probabilístico	Estratificado	Ensayos clínicos	Mayor complejidad
No probabilístico	Conveniencia	Estudios clínicos iniciales	Baja generalización

Nota. La tabla muestra los principales métodos de muestreo y sus aplicaciones habituales en investigación biomédica y clínica.

1.4.4 Parámetros y estadísticos: conceptos fundamentales

En bioestadística es esencial distinguir entre parámetros y estadísticos. Los parámetros son valores numéricos que describen características de una población completa, mientras que los estadísticos son medidas calculadas a partir de una muestra y se utilizan para estimar dichos parámetros.

Por ejemplo, la media poblacional de la presión arterial es un parámetro, generalmente desconocido, mientras que la media calculada en una muestra de pacientes constituye un estadístico. La inferencia estadística se basa en el uso

de estadísticos muestrales para aproximar los parámetros poblacionales con un nivel conocido de incertidumbre (Rosner, 2016).

Tabla 1-10. Relación entre parámetros poblacionales y estadísticos muestrales

Característica	Parámetro	Estadístico
Media	μ	$\bar{x}$
Varianza	σ^2	s^2
Proporción	p	$\hat{p}$

Nota. La tabla muestra la correspondencia entre medidas poblacionales teóricas y estimaciones obtenidas a partir de muestras.

1.4.5 Fundamento matemático de la inferencia: estimación de parámetros

La inferencia estadística se apoya en el principio de que los estadísticos muestrales varían de muestra a muestra, pero lo hacen siguiendo patrones probabilísticos conocidos. Un ejemplo básico es la estimación de la media poblacional mediante la media muestral:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

La variabilidad de esta estimación se cuantifica mediante el error estándar:

$$SE = \frac{s}{\sqrt{n}}$$

Estas expresiones muestran que, a mayor tamaño de muestra, menor es la variabilidad de la estimación, lo que se traduce en mayor precisión. Este principio justifica el cálculo del tamaño muestral antes de iniciar un estudio clínico o epidemiológico.

Ejercicio aplicado 1. Identificación de población, muestra y parámetros

Planteamiento:

Un hospital desea estimar la media de colesterol total en adultos atendidos en consulta externa durante un año. Se selecciona una muestra de 120 pacientes.

Solución:

- Población objetivo: adultos atendidos en consulta externa del hospital
- Muestra: 120 pacientes seleccionados
- Parámetro de interés: media poblacional del colesterol total
- Estadístico utilizado: media muestral del colesterol

Interpretación:

La media calculada en la muestra permite estimar el valor poblacional con un nivel de incertidumbre cuantificable.

Ejercicio aplicado 2. Efecto del tamaño muestral en la precisión**Planteamiento:**

Dos estudios estiman la media de glucosa en ayunas. El estudio A incluye 30 sujetos y el estudio B incluye 120 sujetos, ambos con desviación estándar similar.

Análisis:

El error estándar del estudio B será aproximadamente la mitad del del estudio A, dado que:

$$SE \propto \frac{1}{\sqrt{n}}$$

Conclusión:

El estudio con mayor tamaño muestral produce una estimación más precisa de la media poblacional, aun cuando ambos estudios utilicen el mismo instrumento de medición.

1.5 Fuentes de datos en salud: clínicos, epidemiológicos y experimentales

La calidad de la investigación en ciencias de la salud depende, en gran medida, del origen, la integridad y la pertinencia de los datos utilizados. La bioestadística aporta criterios para seleccionar, organizar y analizar datos provenientes de distintas fuentes, cada una con características propias, alcances específicos y limitaciones metodológicas. En términos generales, los datos

sanitarios se agrupan en tres grandes categorías: datos clínicos, datos epidemiológicos y datos experimentales. Esta clasificación permite orientar el diseño del estudio, el control de calidad de la información y la selección de métodos estadísticos adecuados.

El conocimiento detallado de estas fuentes es esencial para profesionales de la medicina, enfermería, laboratorio clínico y salud pública, ya que una interpretación incorrecta del origen de los datos puede conducir a conclusiones inapropiadas o a decisiones clínicas mal fundamentadas.

1.5.1 Datos clínicos

Los datos clínicos se originan en la atención directa de pacientes y reflejan procesos diagnósticos, terapéuticos y de seguimiento. Constituyen una de las fuentes más utilizadas en investigación biomédica aplicada, debido a su cercanía con la práctica asistencial (Altman, 1991).

Entre los principales datos clínicos se incluyen:

- Historia clínica
- Resultados de laboratorio
- Signos vitales
- Diagnósticos médicos
- Procedimientos terapéuticos
- Evolución clínica y desenlaces

Estos datos pueden registrarse de forma prospectiva o retrospectiva. Los estudios prospectivos planifican la recolección de datos desde el inicio, mientras que los estudios retrospectivos analizan información previamente registrada. Cada enfoque tiene implicaciones distintas en cuanto a control de calidad, sesgos y validez interna.

Un aspecto crítico de los datos clínicos es su heterogeneidad. Las mediciones pueden provenir de distintos profesionales, instrumentos o momentos del proceso asistencial, lo que introduce variabilidad adicional. La bioestadística

permite cuantificar esa variabilidad y evaluar su impacto sobre las conclusiones del estudio (Rosner, 2016).

Tabla 1-11. Características de los datos clínicos en investigación sanitaria

Aspecto	Descripción	Implicación estadística
Origen	Atención directa al paciente	Alta relevancia clínica
Registro	Prospectivo o retrospectivo	Diferencias en control de calidad
Variabilidad	Biológica y de medición	Requiere análisis cuidadoso

Nota. Rasgos esenciales de los datos clínicos y su impacto en el análisis estadístico aplicado a la práctica sanitaria.

1.5.2 Datos epidemiológicos

Los datos epidemiológicos se generan a partir de la observación sistemática de eventos de salud en poblaciones o comunidades. Su objetivo principal es describir la distribución de enfermedades, identificar factores asociados y apoyar la planificación de intervenciones sanitarias (Gordis, 2014).

Este tipo de datos suele obtenerse mediante:

- Sistemas de vigilancia sanitaria
- Registros poblacionales
- Encuestas de salud
- Censos y estadísticas vitales

A diferencia de los datos clínicos, los datos epidemiológicos se enfocan en grupos poblacionales y no en individuos aislados. Por ello, los análisis estadísticos se orientan a estimar tasas, proporciones, razones y medidas de asociación, como incidencia, prevalencia, riesgo relativo u odds ratio (Rothman et al., 2014).

La calidad de los datos epidemiológicos depende de la cobertura, la oportunidad del registro y la definición estandarizada de los eventos. La bioestadística permite evaluar la consistencia de las series temporales, detectar

valores atípicos y analizar tendencias, aspectos fundamentales para la vigilancia de enfermedades y la evaluación de políticas de salud.

Tabla 1-12. Principales fuentes de datos epidemiológicos

Fuente	Tipo de información	Uso principal
Vigilancia sanitaria	Casos notificados	Seguimiento de enfermedades
Encuestas poblacionales	Factores de riesgo	Planificación en salud
Registros vitales	Nacimientos y defunciones	Indicadores sanitarios

Nota. Fuentes habituales de datos poblacionales utilizadas para describir y analizar eventos de salud en comunidades.

1.5.3 Datos experimentales

Los datos experimentales se obtienen cuando el investigador interviene activamente sobre una variable, controlando las condiciones del estudio. En ciencias de la salud, este enfoque se aplica en ensayos clínicos, estudios de laboratorio y experimentos con modelos biológicos (Fisher, 1971; Schulz et al., 2010).

Las características distintivas de los datos experimentales incluyen:

- Asignación controlada de tratamientos o intervenciones
- Comparación entre grupos
- Control sistemático de variables externas

En los ensayos clínicos, los datos experimentales permiten evaluar la eficacia y seguridad de intervenciones terapéuticas. En el laboratorio clínico y biomédico, se emplean para validar métodos analíticos, estudiar mecanismos fisiopatológicos y comparar técnicas diagnósticas.

Desde el punto de vista estadístico, estos datos ofrecen mayores posibilidades de inferencia causal, siempre que se respeten principios metodológicos como

la asignación aleatoria y el control de sesgos. La bioestadística proporciona herramientas para el análisis comparativo, la estimación de efectos y la evaluación de precisión.

Tabla 1-13. Comparación entre fuentes de datos en salud

Fuente	Nivel de control	Alcance	Limitación frecuente
Clínicos	Moderado	Atención individual	Registros incompletos
Epidemiológicos	Bajo	Poblacional	Subregistro
Experimentales	Alto	Evaluación causal	Costo y complejidad

Nota. Tabla de comparación general de fuentes de datos sanitarios según control, alcance y limitaciones metodológicas.

1.5.4 Integración de fuentes de datos y análisis estadístico

En la práctica actual, muchos estudios en salud combinan más de una fuente de datos. Por ejemplo, un ensayo clínico puede complementarse con información epidemiológica para contextualizar resultados, o los datos clínicos pueden integrarse con registros poblacionales para evaluar impacto a largo plazo.

La integración de fuentes requiere procedimientos estadísticos cuidadosos, como armonización de variables, evaluación de consistencia y análisis estratificado. La bioestadística facilita esta integración mediante técnicas de estandarización y modelado, siempre considerando la naturaleza de cada fuente y sus posibles sesgos.

1.6 Errores de medición y sesgos en los estudios de salud

La investigación en ciencias de la salud se desarrolla en escenarios donde la medición de fenómenos biológicos y clínicos está expuesta a variabilidad y a imperfecciones del proceso observacional. La bioestadística aporta herramientas conceptuales y analíticas para identificar, cuantificar y mitigar errores de medición y sesgos, con el objetivo de proteger la validez de los resultados y la interpretación clínica. La metodología en estadística médica y

epidemiología ha señalado de forma reiterada que una proporción relevante de conclusiones incorrectas se origina en fallas de medición o en distorsiones sistemáticas durante el diseño y la ejecución de los estudios (Altman, 1991; Rothman et al., 2008).



Figura 1-5. Fuentes de error y sesgo en estudios de salud

1.6.1 Error de medición en ciencias de la salud

El error de medición se define como la diferencia entre el valor observado y el valor verdadero de una variable. En salud, dicho valor verdadero suele ser desconocido, por lo que el análisis se centra en la naturaleza y el comportamiento del error. Textos clásicos de bioestadística distinguen dos tipos principales de error de medición: error aleatorio y error sistemático.

a) Error aleatorio

El error aleatorio se produce por fluctuaciones impredecibles asociadas al proceso de medición. Puede deberse a variaciones biológicas normales, imprecisión del instrumento o condiciones ambientales. Este tipo de error tiende a dispersar las mediciones alrededor del valor verdadero sin desplazarlo de forma consistente.

En el laboratorio clínico, el error aleatorio se manifiesta como imprecisión analítica y suele evaluarse mediante la desviación estándar o el coeficiente de variación. Desde el punto de vista estadístico, el aumento del tamaño muestral reduce el impacto del error aleatorio sobre las estimaciones.

b) Error sistemático

El error sistemático ocurre cuando las mediciones se desvían de forma consistente en una misma dirección respecto al valor verdadero. Puede originarse por instrumentos mal calibrados, procedimientos inadecuados o errores en la definición operacional de la variable (Altman, 1991).

A diferencia del error aleatorio, el error sistemático no se corrige aumentando el tamaño muestral. Su presencia compromete la validez interna del estudio y puede conducir a conclusiones clínicas erróneas, incluso con análisis estadísticos correctos.

Tabla 1-14. Tipos de error de medición y efectos en estudios de salud

Tipo de error	Característica	Ejemplo sanitario	Consecuencia analítica
Aleatorio	Variación impredecible	Variación biológica	Menor precisión
Sistemático	Desplazamiento constante	Equipo mal calibrado	Estimación sesgada

Nota. Diferencias entre errores de medición y su impacto sobre precisión y validez de resultados en investigación sanitaria.

1.6.2 Sesgos en investigación biomédica y clínica

El sesgo se refiere a una distorsión sistemática en la estimación del efecto o asociación de interés, producto del diseño, la recolección de datos, el análisis o la interpretación. A diferencia del error aleatorio, el sesgo altera la dirección o magnitud del resultado de manera consistente.

a) Sesgo de selección

El sesgo de selección aparece cuando la forma de inclusión de los participantes genera grupos no comparables. En estudios clínicos, puede presentarse si los pacientes más graves o más sanos tienen mayor probabilidad de ser incluidos. En estudios observacionales, este sesgo limita la generalización de los resultados.

La asignación aleatoria en ensayos clínicos y la definición clara de criterios de inclusión y exclusión son estrategias reconocidas para reducir este tipo de sesgo.

b) Sesgo de información

El sesgo de información surge cuando la medición de la exposición o del desenlace difiere entre grupos. Ejemplos frecuentes incluyen errores de clasificación diagnóstica, registros incompletos y diferencias en la forma de recolectar datos clínicos.

En estudios retrospectivos, la calidad del registro clínico es un factor determinante para minimizar este sesgo.

c) Sesgo de confusión

La confusión ocurre cuando una tercera variable se asocia tanto con la exposición como con el desenlace, generando una asociación aparente que no corresponde a una relación directa. Manuales de epidemiología señalan que la confusión es una amenaza frecuente en estudios observacionales y requiere control mediante diseño o análisis estadístico.

Tabla 1-15. Sesgos frecuentes en estudios de salud

Sesgo	Origen	Ejemplo	Estrategia de control
Selección	Inclusión no comparable	Estudios hospitalarios	Aleatorización
Información	Medición diferencial	Registro clínico	Estandarización

Nota. Sesgos habituales que afectan la validez de estudios clínicos y epidemiológicos y formas generales de mitigación.

1.6.3 Evaluación estadística del error y control metodológico

La bioestadística ofrece procedimientos para evaluar el impacto del error de medición. En variables continuas, la variabilidad se cuantifica mediante la varianza y la desviación estándar:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

En el laboratorio clínico, el coeficiente de variación se utiliza para expresar imprecisión relativa:

$$CV(\%) = \frac{s}{\bar{x}} \times 100$$

Estos indicadores permiten comparar la precisión de métodos y decidir su idoneidad para uso clínico o investigativo (Motulsky, 2015).

1.7 Ética y confidencialidad en el manejo de datos sanitarios

El manejo de datos en ciencias de la salud implica responsabilidades éticas y legales que trascienden el análisis estadístico. Los datos sanitarios contienen información sensible que puede afectar la privacidad, la dignidad y los derechos de las personas. Por ello, la bioestadística aplicada se integra con principios éticos que orientan la recolección, el almacenamiento, el análisis y la difusión de la información.

1.7.1 Principios éticos aplicables a los datos en salud

La ética de la investigación biomédica se sustenta en principios ampliamente aceptados, como el respeto por las personas, la beneficencia y la justicia. Estos principios se reflejan en prácticas concretas de manejo de datos, como el consentimiento informado y la protección de la confidencialidad.

El consentimiento informado implica que los participantes conozcan el uso previsto de sus datos, los riesgos asociados y las medidas de protección

adoptadas. Desde la perspectiva estadística, esto incluye informar sobre la finalidad del análisis y la forma de presentación de resultados agregados.

1.7.2 Confidencialidad y anonimización de datos

La confidencialidad se refiere a la obligación de proteger la información personal de accesos no autorizados. En investigación sanitaria, se emplean técnicas de anonimización o seudonimización para reducir el riesgo de identificación de los participantes.

La bioestadística favorece el uso de bases de datos anonimizadas, donde los identificadores directos han sido eliminados o codificados, permitiendo el análisis sin comprometer la identidad de los sujetos.

Tabla 1-16. Medidas de protección de datos sanitarios

Medida	Descripción	Aplicación
Anonimización	Eliminación de identificadores	Bases de datos
Seudonimización	Reemplazo por códigos	Seguimiento
Acceso restringido	Control de usuarios	Seguridad

Nota. La tabla muestra las estrategias básicas para proteger la privacidad de participantes en investigaciones con datos de salud.

1.7.3 Ética del análisis y comunicación de resultados

La ética no se limita a la recolección de datos. También abarca el análisis y la comunicación de los resultados. Prácticas como la manipulación selectiva de datos, la omisión de resultados no favorables o la interpretación exagerada de hallazgos contravienen los principios éticos de la investigación.

Las guías internacionales de reporte recomiendan presentar métodos y resultados de forma completa y transparente, permitiendo la evaluación crítica por parte de la comunidad científica.

Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

CAPÍTULO II

Organización y Descripción de Datos Biomédicos

AUTOR: Azúa Menéndez Marieta Del Jesús



2 Organización y Descripción de Datos Biomédicos

2.1 Tabulación y clasificación de datos clínicos

La organización inicial de los datos constituye una etapa decisiva en la investigación biomédica y clínica. Antes de aplicar cualquier técnica descriptiva o inferencial, es imprescindible ordenar la información recolectada de forma sistemática, coherente y verificable. La tabulación y clasificación de datos clínicos permiten transformar registros dispersos en estructuras comprensibles que facilitan la exploración de patrones, la detección de inconsistencias y la preparación del análisis estadístico posterior (W. Daniel & Cross, 2013; Rosner Bernard, 2016).

En entornos sanitarios, los datos suelen provenir de historias clínicas, formularios de investigación, resultados de laboratorio y sistemas de información hospitalaria. Estos registros pueden contener variables de distinta naturaleza, diferentes unidades de medida y múltiples momentos de observación. La bioestadística aplicada proporciona criterios para clasificar y tabular esta información sin perder significado clínico ni introducir distorsiones analíticas.



Figura 2-1. Proceso de tabulación de datos clínicos

2.1.1 Finalidad de la tabulación en datos clínicos

La tabulación consiste en disponer los datos en filas y columnas siguiendo un orden lógico que refleje la estructura del estudio. En investigación clínica, este procedimiento cumple varias funciones:

- Facilita la revisión de la calidad de los datos y la detección de valores faltantes o inconsistentes.
- Permite identificar la naturaleza de las variables y su escala de medición.
- Sirve como base para la construcción de tablas de frecuencia, gráficos y cálculos estadísticos.

Expertos en estadística médica señalan que una tabulación inadecuada puede inducir errores posteriores, incluso cuando se utilizan métodos analíticos correctos (Altman, 1991; Motulsky, 2015). Por ello, la organización de los datos no debe considerarse una tarea mecánica, sino una fase analítica con implicaciones metodológicas claras.

2.1.2 Clasificación inicial de datos clínicos

La clasificación de datos clínicos implica agrupar la información según criterios previamente definidos, acordes con los objetivos del estudio. Esta clasificación puede realizarse atendiendo a distintos enfoques:

- Según el tipo de variable: cualitativa o cuantitativa.
- Según el momento de medición: basal, seguimiento, desenlace.
- Según el origen del dato: clínico, de laboratorio o administrativo.

Esta etapa se apoya directamente en los conceptos desarrollados en el Capítulo 1 sobre tipos de variables y escalas de medición. Una clasificación incorrecta, por ejemplo, tratar una variable ordinal como continua, compromete la validez del análisis descriptivo posterior (Rosner, 2016).

Tabla 2-1. Ejemplo de clasificación de datos clínicos en un estudio hospitalario

Variable	Tipo de variable	Escala	Origen
Edad	Cuantitativa continua	Razón	Historia clínica
Sexo	Cualitativa nominal	Nominal	Registro administrativo
Glucosa en ayunas	Cuantitativa continua	Razón	Laboratorio
Diagnóstico	Cualitativa nominal	Nominal	Evaluación médica
Grado de dolor	Cualitativa ordinal	Ordinal	Escala clínica

Nota. La tabla muestra la clasificación inicial de variables clínicas según tipo, escala y fuente para facilitar la organización y análisis estadístico.

2.1.3 Estructura básica de una tabla de datos clínicos

Una tabla de datos clínicos suele organizarse de modo que cada fila represente una unidad de observación, generalmente un paciente, y cada columna corresponda a una variable. Esta disposición permite aplicar de forma directa procedimientos estadísticos y facilita la importación de datos a software especializado.

Ejemplo conceptual de estructura tabular:

- Filas: pacientes o registros individuales.
- Columnas: variables clínicas, demográficas o de laboratorio.

Es recomendable utilizar códigos claros y consistentes para las variables cualitativas, evitando ambigüedades. Por ejemplo, codificar sexo como 1 = masculino, 2 = femenino, siempre acompañado de un diccionario de datos que documente estas decisiones.

2.1.4 Control de calidad durante la tabulación

La tabulación ofrece una oportunidad temprana para el control de calidad de los datos. En esta fase se pueden identificar:

- Valores fuera de rango fisiológicamente posibles.

- Inconsistencias lógicas entre variables, como edad pediátrica con diagnóstico exclusivo de adultos.
- Registros incompletos o duplicados.

La bioestadística recomienda aplicar reglas simples de validación antes de avanzar al análisis. Por ejemplo, verificar que la presión arterial sistólica se encuentre dentro de un rango plausible o que las fechas sigan una secuencia temporal coherente.

2.1.5 Preparación matemática básica de los datos

En esta etapa inicial, no se realizan aún análisis complejos, pero sí se aplican operaciones básicas que permiten verificar la coherencia numérica de los datos cuantitativos. Por ejemplo, el cálculo preliminar de la media puede ayudar a detectar valores atípicos evidentes:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

Si un valor individual se encuentra muy alejado del resto, se justifica revisar el registro original antes de continuar. Esta práctica preventiva es ampliamente recomendada en textos de bioestadística aplicada (Rosner, 2016).

2.1.6 Importancia clínica de una correcta tabulación

Desde la perspectiva clínica, la tabulación adecuada de los datos garantiza que los resultados estadísticos reflejen de manera fiel la realidad observada en los pacientes. Una mala organización puede ocultar tendencias relevantes, exagerar diferencias o generar interpretaciones erróneas con impacto directo en la toma de decisiones sanitarias.

Además, una base de datos bien estructurada facilita la reproducibilidad del estudio y la auditoría metodológica, aspectos cada vez más valorados en la investigación biomédica contemporánea.

2.1.6.1 *Ejercicio: tabulación y clasificación de datos clínicos*

Planteamiento

En un estudio descriptivo se recolectan los siguientes datos de 10 pacientes atendidos en consulta externa:

- Edad (años)
- Sexo
- Presión arterial sistólica (mmHg)
- Diagnóstico de hipertensión
- Número de consultas en el último año

Actividad

1. Clasificar cada variable según su tipo y escala de medición.
2. Organizar los datos clínicos en una tabla adecuada para su análisis estadístico.

Solución

Tabla 2-2. Clasificación de las variables

Variable	Tipo de variable	Escala de medición
Edad	Cuantitativa continua	Razón
Sexo	Cualitativa nominal	Nominal
Presión arterial sistólica	Cuantitativa continua	Razón
Diagnóstico de hipertensión	Cualitativa nominal dicotómica	Nominal
Número de consultas	Cuantitativa discreta	Razón

Tabulación de los datos clínicos

Para la organización de los datos se adopta la siguiente codificación:

- Sexo: 1 = Masculino, 2 = Femenino
- Diagnóstico de hipertensión: 0 = No, 1 = Sí

Tabla 2-3. Base de datos clínica organizada para análisis estadístico

Paciente	Edad (años)	Sexo	Presión sistólica (mmHg)	Hipertensión	Consultas/año
1	45	1	138	1	4
2	52	2	150	1	6
3	39	1	120	0	2
4	61	2	160	1	8
5	48	1	130	0	3
6	55	2	145	1	5
7	42	1	125	0	1
8	67	2	170	1	9
9	50	1	140	1	4
10	36	2	118	0	2

Nota. La tabla muestra un ejemplo de tabulación clínica con codificación estándar y estructura compatible con análisis estadístico descriptivo.

Interpretación

La tabla organizada presenta una fila por paciente y una columna por variable, con codificación clara de las variables cualitativas. Esta estructura permite aplicar de forma directa tablas de frecuencia, gráficos y medidas descriptivas, así como importar los datos sin modificaciones a software estadístico. Una tabulación correcta en esta etapa reduce errores posteriores y asegura coherencia metodológica en el análisis biomédico.

2.2 Tablas de frecuencia y representación de resultados de laboratorio

La descripción inicial de los datos biomédicos requiere herramientas que permitan resumir grandes volúmenes de información sin perder significado clínico. Entre estas herramientas, las tablas de frecuencia ocupan un lugar central, ya que facilitan la visualización ordenada de la distribución de los datos y constituyen la base para la interpretación descriptiva y la posterior representación gráfica. En el ámbito de la salud, las tablas de frecuencia se utilizan de forma habitual para resumir resultados de laboratorio, variables clínicas categóricas y mediciones cuantitativas agrupadas (Rosner, 2016; Daniel y Cross, 2017).

En estudios clínicos y de laboratorio, una tabla de frecuencia bien construida permite identificar patrones, detectar valores inusuales y comparar resultados con rangos de referencia. Por esta razón, su elaboración debe respetar criterios estadísticos y clínicos, evitando agrupaciones arbitrarias o interpretaciones alejadas de la realidad biomédica.

2.2.1 Concepto y finalidad de las tablas de frecuencia

Una tabla de frecuencia es una organización sistemática de los datos en categorías o intervalos, acompañada del número de observaciones que corresponden a cada uno. Su finalidad principal es resumir la información, facilitando la comprensión de la distribución de una variable (Motulsky, 2015).

En bioestadística aplicada a la salud, las tablas de frecuencia permiten:

- Describir la distribución de variables cualitativas y cuantitativas.
- Comparar resultados de laboratorio entre grupos.
- Identificar concentraciones de valores dentro o fuera de rangos clínicos esperados.

La elección del tipo de tabla depende del tipo de variable y del objetivo del análisis descriptivo.

2.2.2 Tablas de frecuencia para variables cualitativas

En las variables cualitativas, las tablas de frecuencia presentan el número de casos en cada categoría. En investigación clínica, estas tablas son comunes para describir sexo, diagnósticos, presencia o ausencia de una condición o resultados categóricos de pruebas diagnósticas (Altman, 1991).

Por ejemplo, al analizar el diagnóstico de hipertensión en un grupo de pacientes, se emplea una tabla de frecuencia simple que muestre el número y porcentaje de casos con y sin la condición.

Tabla 2-4. Tabla de frecuencia del diagnóstico de hipertensión

Diagnóstico	Frecuencia	Porcentaje (%)
No	4	40.0
Sí	6	60.0
Total	10	100.0

Nota. Distribución de una variable cualitativa dicotómica expresada mediante frecuencias absolutas y porcentajes.

Este tipo de tabla permite una lectura inmediata del predominio de la condición estudiada y resulta esencial en la descripción inicial de muestras clínicas.

2.2.3 Tablas de frecuencia para variables cuantitativas discretas

Las variables cuantitativas discretas, como el número de consultas médicas o episodios de enfermedad, también pueden resumirse mediante tablas de frecuencia. En este caso, cada valor observado se asocia con el número de veces que aparece en la muestra (Daniel y Cross, 2017).

Tabla 2-5. Tabla de frecuencia del número de consultas médicas anuales

Consultas/año	Frecuencia
1	1
2	2

3	1
4	2
5	1
6	1
8	1
9	1
Total	10

Nota. Distribución de una variable cuantitativa discreta sin agrupación en intervalos.

Esta forma de tabulación es adecuada cuando el número de valores distintos es reducido y clínicamente interpretable.

2.2.4 Tablas de frecuencia agrupadas para variables cuantitativas continuas

Cuando se trabaja con variables cuantitativas continuas, como concentraciones bioquímicas o valores hematológicos, suele ser necesario agrupar los datos en intervalos o clases. Esta práctica facilita la lectura y permite describir la forma general de la distribución.

En resultados de laboratorio, la agrupación debe considerar rangos clínicos relevantes y la amplitud de los valores observados. Un criterio común consiste en definir intervalos de igual amplitud, procurando que el número de clases no sea excesivo ni insuficiente.

Tabla 2-6. Tabla de frecuencia agrupada de glucosa en ayunas

Intervalo (mg/dL)	Frecuencia
110 – 119	1
120 – 129	2
130 – 139	2
140 – 149	2

150 – 159	1
160 – 169	1
170 – 179	1
Total	10

Nota. Agrupación de valores continuos de laboratorio en intervalos clínicamente interpretables.

Este tipo de tabla permite identificar concentraciones frecuentes y facilita la comparación con valores de referencia establecidos en el laboratorio clínico.

2.2.5 Frecuencia relativa y frecuencia acumulada

Además de la frecuencia absoluta, es habitual calcular la frecuencia relativa y la frecuencia acumulada, especialmente en variables cuantitativas. La frecuencia relativa expresa la proporción de observaciones en cada categoría o intervalo, mientras que la frecuencia acumulada indica el número total de observaciones hasta un determinado punto de la distribución.

La frecuencia relativa se calcula mediante la expresión:

$$f_r = \frac{f_i}{n}$$

donde f_i es la frecuencia absoluta de la clase y n el tamaño total de la muestra.

Tabla 2-7. Frecuencias absoluta, relativa y acumulada de glucosa en ayunas

Intervalo (mg/dL)	f	fr	Fa
110 – 119	1	0.10	1
120 – 129	2	0.20	3
130 – 139	2	0.20	5
140 – 149	2	0.20	7
150 – 159	1	0.10	8

160 – 169	1	0.10	9
170 – 179	1	0.10	10

Nota. Ejemplo de cálculo de frecuencias absoluta, relativa y acumulada en resultados de laboratorio.

2.2.6 Representación de resultados de laboratorio y utilidad clínica

La representación tabular de resultados de laboratorio facilita la comparación con rangos de referencia y la identificación de posibles alteraciones clínicas. En estudios descriptivos, estas tablas permiten evaluar la distribución general de los valores y detectar concentraciones de resultados patológicos o limítrofes.

En la práctica clínica, una correcta tabulación favorece la comunicación entre profesionales de la salud y constituye un soporte objetivo para la toma de decisiones diagnósticas y terapéuticas.

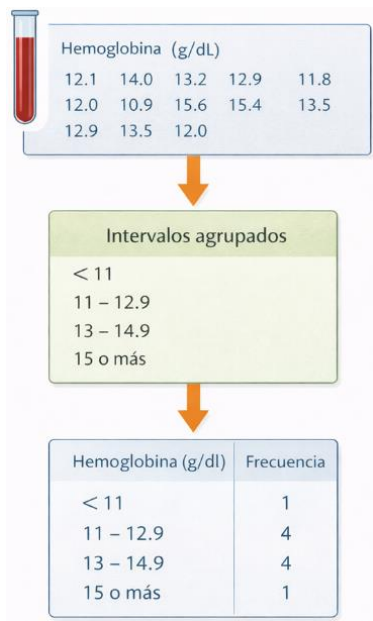


Figura 2-2. Relación entre datos crudos y tablas de frecuencia

2.2.6.1 Ejercicio: construcción e interpretación de una tabla de frecuencia

Planteamiento

Se dispone de los valores de colesterol total (mg/dL) de 12 pacientes: 185, 192, 198, 205, 210, 218, 225, 232, 240, 245, 255, 268.

Actividad

1. Construir una tabla de frecuencia agrupada con intervalos de 20 mg/dL.
2. Calcular la frecuencia absoluta y la frecuencia relativa.
3. Interpretar los resultados desde el punto de vista clínico.

Solución

Tabla 2-8. Tabla de frecuencia del colesterol total

Intervalo (mg/dL)	Frecuencia	Frecuencia relativa
180 – 199	3	0.25
200 – 219	3	0.25
220 – 239	3	0.25
240 – 259	2	0.17
260 – 279	1	0.08
Total	12	1.00

Nota. La tabla muestra la distribución agrupada de colesterol total para descripción clínica y comparación con rangos de referencia.

Interpretación

La mayor concentración de valores se observa entre 180 y 239 mg/dL, lo que indica predominio de niveles cercanos al límite superior recomendado. La tabla permite identificar pacientes con valores elevados que podrían requerir evaluación clínica adicional.

2.3 Diagramas de barras, histogramas y polígonos de frecuencia

La representación gráfica constituye un complemento indispensable de las tablas de frecuencia en la descripción de datos biomédicos. En ciencias de la salud, los gráficos permiten identificar patrones, asimetrías, concentraciones de valores y posibles valores atípicos con mayor rapidez que la inspección tabular. La elección del tipo de diagrama debe responder a la naturaleza de la variable, a su escala de medición y al objetivo clínico del análisis descriptivo.

Entre los gráficos más utilizados en la práctica biomédica se encuentran los diagramas de barras, los histogramas y los polígonos de frecuencia. Cada uno cumple una función específica y presenta ventajas particulares cuando se emplea de forma correcta.

2.3.1 Diagramas de barras

El diagrama de barras se utiliza principalmente para representar variables cualitativas y cuantitativas discretas. Consiste en barras rectangulares separadas entre sí, cuya altura es proporcional a la frecuencia absoluta o relativa de cada categoría. En el ámbito clínico, este tipo de gráfico es común para describir diagnósticos, resultados categóricos de pruebas, sexo, grupos etarios y estados clínicos.

Una característica esencial del diagrama de barras es que las categorías no mantienen continuidad numérica, por lo que las barras deben aparecer separadas. La bioestadística aplicada recomienda ordenar las categorías de forma lógica o clínica para facilitar la interpretación, por ejemplo, ordenar diagnósticos por frecuencia o gravedad.

Tabla 2-9. Distribución de pacientes según diagnóstico de hipertensión

Diagnóstico	Frecuencia
No	4
Sí	6
Total	10

Nota. La tabla muestra frecuencias de una variable cualitativa dicotómica utilizadas como base para un diagrama de barras.

A partir de esta tabla, el diagrama de barras permite visualizar de forma inmediata el predominio del diagnóstico positivo en la muestra.

2.3.2 Histogramas

El histograma es el gráfico de elección para variables cuantitativas continuas. Representa la distribución de los datos mediante barras contiguas, cada una correspondiente a un intervalo de clase. A diferencia del diagrama de barras, las barras no están separadas, lo que refleja la continuidad de la variable (Pagano y Gauvreau, 2018).

En resultados de laboratorio, los histogramas son especialmente útiles para evaluar la forma de la distribución, identificar asimetrías y valorar la proximidad de los valores a rangos de referencia. Textos de estadística médica señalan que el histograma es una herramienta inicial para decidir la aplicabilidad de métodos descriptivos y analíticos posteriores (Rosner, 2016).

Tabla 2-10. Intervalos de glucosa en ayunas y frecuencias

Intervalo (mg/dL)	Frecuencia
110 – 119	1
120 – 129	2
130 – 139	2
140 – 149	2
150 – 159	1
160 – 169	1
170 – 179	1

Nota. Tabla de frecuencias agrupadas utilizada para la construcción de un histograma de laboratorio.

En el histograma correspondiente, cada barra representa un intervalo de glucosa, y su área es proporcional a la frecuencia observada.

2.3.3 Polígonos de frecuencia

El polígono de frecuencia es una representación gráfica que se construye uniendo mediante segmentos de línea los puntos medios de los intervalos de clase, con alturas proporcionales a las frecuencias. Este gráfico se utiliza con variables cuantitativas continuas y resulta útil para comparar distribuciones entre grupos o para observar la tendencia general de los datos.

En estudios biomédicos, el polígono de frecuencia permite superponer distribuciones de diferentes grupos, por ejemplo, comparar valores de laboratorio antes y después de una intervención, manteniendo claridad visual.

El punto medio de cada intervalo se calcula mediante la expresión:

$$x_m = \frac{L_i + L_s}{2}$$

donde L_i es el límite inferior del intervalo y L_s el límite superior.

Tabla 2-11. Puntos medios para la construcción del polígono de frecuencia

Intervalo (mg/dL)	Punto medio	Frecuencia
110 – 119	114.5	1
120 – 129	124.5	2
130 – 139	134.5	2
140 – 149	144.5	2
150 – 159	154.5	1
160 – 169	164.5	1
170 – 179	174.5	1

Nota. Cálculo de puntos medios para representar gráficamente la distribución mediante un polígono de frecuencia.

2.3.4 Comparación y criterios de selección del gráfico

La selección adecuada del gráfico depende del tipo de variable y del objetivo del análisis descriptivo. Utilizar un gráfico inapropiado puede inducir interpretaciones erróneas, aun cuando los datos estén correctamente tabulados.

Tabla 2-12. Criterios para seleccionar el tipo de gráfico en datos biomédicos

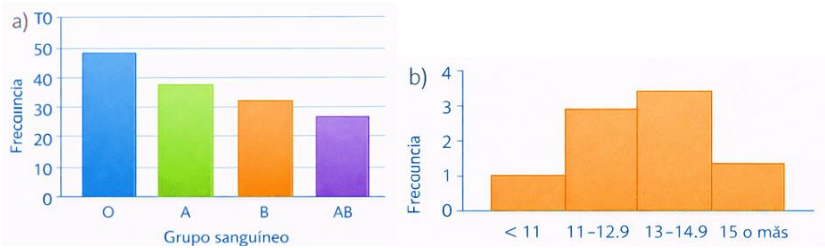
Tipo de variable	Gráfico recomendado	Uso principal
Cualitativa	Diagrama de barras	Comparar categorías
Cuantitativa continua	Histograma	Evaluar distribución
Cuantitativa continua	Polígono de frecuencia	Comparar tendencias

Nota. Relación entre tipo de variable biomédica y representación gráfica recomendada.

2.3.5 Interpretación clínica de los gráficos

La interpretación gráfica debe realizarse siempre en conjunto con el conocimiento clínico. Un histograma de laboratorio con valores desplazados hacia intervalos altos puede sugerir alteraciones metabólicas, mientras que un diagrama de barras con predominio de una categoría puede orientar decisiones clínicas o administrativas.

La bioestadística aplicada enfatiza que los gráficos no sustituyen el análisis numérico, sino que lo complementan, aportando una visión global que orienta la interpretación y la comunicación de resultados.



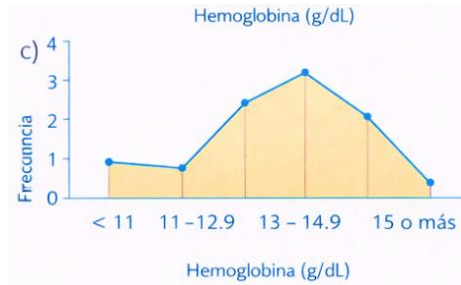


Figura 2-3. Representaciones gráficas de datos biomédicos

2.3.5.1 Ejercicio propuesto: construcción e interpretación de gráficos

Planteamiento

Se dispone de los valores de presión arterial sistólica (mmHg) de 12 pacientes: 118, 122, 130, 135, 138, 140, 145, 150, 155, 160, 165, 170.

Actividad

1. Construir una tabla de frecuencia agrupada con intervalos de 10 mmHg.
2. Elaborar el histograma y el polígono de frecuencia.
3. Interpretar clínicamente la distribución observada.

Solución

Tabla 2-13. Frecuencia de presión arterial sistólica

Intervalo (mmHg)	Frecuencia
110 – 119	1
120 – 129	1
130 – 139	3
140 – 149	2
150 – 159	2

160 – 169	2
170 – 179	1
Total	12

Nota. Tabla base para la construcción del histograma y polígono de frecuencia de presión arterial sistólica.

Interpretación:

El histograma muestra concentración de valores entre 130 y 159 mmHg, lo que indica un predominio de cifras elevadas. El polígono de frecuencia evidencia una tendencia ascendente hacia valores superiores, compatible con una población con alta proporción de pacientes hipertensos.

2.4 Diagramas de dispersión y correlación visual

El análisis descriptivo en ciencias de la salud no se limita a estudiar variables de manera aislada. En numerosos entornos clínicos y biomédicos resulta necesario explorar la relación entre dos variables cuantitativas, con el fin de identificar tendencias, asociaciones y posibles patrones de dependencia. El diagrama de dispersión constituye la herramienta gráfica básica para este propósito, mientras que la correlación visual permite una primera apreciación de la intensidad y dirección de la relación antes de aplicar métodos analíticos formales (Rosner, 2016; Daniel y Cross, 2017).

En la práctica clínica y de laboratorio, este tipo de representación es frecuente para analizar, por ejemplo, la relación entre edad y presión arterial, dosis y respuesta farmacológica, índice de masa corporal y glucosa, o tiempo de evolución y marcadores bioquímicos. Su correcta utilización facilita la interpretación clínica preliminar y orienta la selección de modelos estadísticos posteriores (Motulsky, 2015).

2.4.1 Concepto y estructura del diagrama de dispersión

Un diagrama de dispersión es una representación gráfica en la que cada observación se muestra como un punto en un plano cartesiano. El eje horizontal suele corresponder a la variable independiente o explicativa, mientras que el eje vertical representa la variable dependiente o de respuesta. Cada punto del

gráfico corresponde a un par de valores observados en un mismo individuo o unidad de estudio.

Este tipo de gráfico es adecuado exclusivamente para variables cuantitativas medidas en escala de intervalo o de razón. Su principal utilidad radica en mostrar de forma directa la forma de la relación entre variables, sin imponer supuestos matemáticos previos.

2.4.2 Interpretación visual de la relación entre variables

La correlación visual se refiere a la evaluación cualitativa de la relación entre dos variables a partir del patrón que adoptan los puntos en el diagrama de dispersión. En términos generales, pueden identificarse tres situaciones básicas:

- Relación positiva, cuando los valores de ambas variables tienden a aumentar conjuntamente.
- Relación negativa, cuando el aumento de una variable se asocia con la disminución de la otra.
- Ausencia de relación aparente, cuando los puntos se distribuyen sin un patrón reconocible.

La estadística médica advierte que la correlación visual no sustituye a la medición cuantitativa de la asociación, pero constituye un paso inicial indispensable para detectar relaciones no lineales, valores atípicos o agrupamientos inesperados.

Tabla 2-14. Ejemplo de pares de datos clínicos para análisis gráfico

Paciente	Edad (años)	Presión sistólica (mmHg)
1	35	118
2	42	125
3	48	132
4	50	138

5	55	145
6	60	150
7	65	158
8	70	165

Nota. Pares de variables cuantitativas utilizados para ilustrar un diagrama de dispersión clínicamente interpretable.

A partir de estos datos, el diagrama de dispersión permite observar una tendencia creciente entre edad y presión arterial sistólica.

2.4.3 Identificación de patrones y valores atípicos

Una ventaja fundamental del diagrama de dispersión es la detección temprana de valores atípicos. Un punto que se aleja claramente del patrón general puede indicar un error de medición, una condición clínica particular o una subpoblación con características diferentes. En estudios biomédicos, estos valores deben revisarse con cautela antes de decidir su exclusión, ya que pueden contener información clínicamente relevante (Rosner, 2016).

Asimismo, el diagrama permite identificar relaciones no lineales, como patrones curvilíneos o agrupamientos, que no serían evidentes a partir de medidas resumen simples.

2.4.4 Relación entre correlación visual y correlación cuantitativa

Aunque la correlación visual aporta una apreciación inicial, su interpretación debe complementarse con medidas numéricas de asociación. En bioestadística, la correlación lineal entre dos variables cuantitativas se cuantifica habitualmente mediante el coeficiente de correlación de Pearson, cuya expresión general es:

$$r = \frac{\sum(x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum(x_i - \bar{x})^2 \sum(y_i - \bar{y})^2}}$$

Este coeficiente toma valores entre -1 y 1 , indicando la dirección y la intensidad de la relación lineal. No obstante, antes de calcularlo, es recomendable inspeccionar el diagrama de dispersión para verificar la forma

de la relación y la presencia de valores atípicos, tal como recomiendan textos clásicos de estadística aplicada a la salud.

Tabla 2-15. Correspondencia entre patrón visual y correlación aproximada

Patrón observado	Tendencia visual	Interpretación general
Puntos alineados crecientes	Relación positiva fuerte	Asociación directa
Puntos alineados decrecientes	Relación negativa fuerte	Asociación inversa
Puntos dispersos	Sin patrón claro	Asociación débil o nula

Nota. Relación orientativa entre patrones gráficos y la intensidad cualitativa de la asociación entre variables.

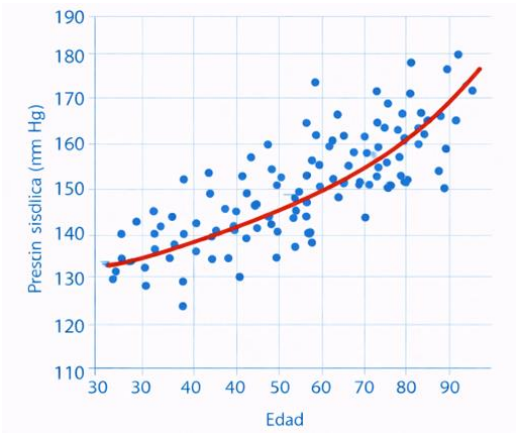


Figura 2-4. Diagrama de dispersión y correlación visual

Ejercicio propuesto: construcción e interpretación de un diagrama de dispersión

Planteamiento

Se registran los valores de índice de masa corporal y glucosa en ayunas (mg/dL) en ocho pacientes:

Paciente	IMC (kg/m ²)	Glucosa (mg/dL)
1	22.5	95
2	24.0	98
3	26.8	105
4	28.5	112
5	30.0	118
6	31.5	125
7	33.0	132
8	35.2	140

Actividad

1. Construir el diagrama de dispersión con IMC en el eje horizontal y glucosa en el eje vertical.
2. Describir la relación observada mediante correlación visual.
3. Indicar si resulta razonable aplicar un análisis de correlación lineal.

Solución

Al representar los pares de valores en un plano cartesiano, los puntos muestran una tendencia creciente clara, lo que indica una relación positiva entre IMC y glucosa en ayunas. No se observan valores atípicos evidentes ni cambios bruscos en la pendiente, por lo que la relación puede considerarse aproximadamente lineal en el rango estudiado.

Desde el punto de vista clínico, esta representación muestra que valores más elevados de IMC se asocian con concentraciones mayores de glucosa, lo que resulta coherente con el conocimiento fisiopatológico. En consecuencia, es razonable complementar el análisis descriptivo con un coeficiente de correlación lineal y, de ser pertinente, con un modelo de regresión simple.

2.5 Medidas de tendencia central: media, mediana y moda

La descripción numérica de los datos biomédicos requiere indicadores que resuman, en un solo valor, la localización o centro de una distribución. Las medidas de tendencia central cumplen esta función y permiten representar el valor típico de una variable clínica o de laboratorio, facilitando la comparación entre grupos y la interpretación de resultados. En ciencias de la salud, las medidas más utilizadas son la media, la mediana y la moda, cada una con propiedades específicas y aplicaciones particulares según el tipo de variable y la forma de la distribución (Rosner, 2016; Daniel y Cross, 2017).

La selección adecuada de la medida de tendencia central es una decisión metodológica relevante. Utilizar una medida inapropiada puede conducir a interpretaciones clínicas erróneas, especialmente en presencia de asimetría o valores atípicos. Textos de estadística médica enfatizan que la medida debe elegirse considerando la naturaleza de los datos y el objetivo del análisis descriptivo (Altman, 1991; Motulsky, 2015).

2.5.1 Media aritmética

La media aritmética es la medida de tendencia central más conocida y utilizada. Se define como la suma de todos los valores observados dividida por el número total de observaciones. En investigación biomédica, la media se emplea con frecuencia para resumir variables cuantitativas continuas como presión arterial, glucosa, colesterol o concentraciones hormonales (Pagano y Gauvreau, 2018).

La expresión matemática de la media es:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

donde x_i representa cada observación y n el tamaño de la muestra.

La media utiliza toda la información disponible, lo que constituye una ventaja desde el punto de vista estadístico. Sin embargo, es sensible a valores extremos. En presencia de valores muy altos o muy bajos, la media puede desplazarse y dejar de representar adecuadamente al conjunto de datos.

2.5.2 Mediana

La mediana es el valor que divide la distribución ordenada en dos partes iguales, de modo que el 50 por ciento de las observaciones se sitúan por debajo y el 50 por ciento por encima. Esta medida resulta especialmente útil cuando los datos presentan asimetría o valores atípicos, situaciones frecuentes en variables clínicas como duración de hospitalización, costos sanitarios o niveles de marcadores biológicos con distribución sesgada (Motulsky, 2015).

Para calcular la mediana, los datos deben ordenarse de menor a mayor. Si el número de observaciones es impar, la mediana corresponde al valor central. Si es par, la mediana se obtiene como el promedio de los dos valores centrales. Desde el punto de vista clínico, la mediana suele reflejar mejor el valor típico de la población cuando existen observaciones extremas que podrían distorsionar la media (Altman, 1991).

2.5.3 Moda

La moda es el valor que aparece con mayor frecuencia en un conjunto de datos. A diferencia de la media y la mediana, la moda puede aplicarse tanto a variables cuantitativas como cualitativas. En ciencias de la salud, la moda se utiliza con frecuencia para describir categorías más comunes, como diagnósticos predominantes, grupos etarios más frecuentes o resultados categóricos de pruebas diagnósticas.

Una distribución puede ser unimodal, bimodal o multimodal, dependiendo del número de valores con máxima frecuencia. La presencia de más de una moda puede indicar heterogeneidad en la población o la coexistencia de subgrupos clínicos con características distintas.

Tabla 2-16. Características de las medidas de tendencia central

Medida	Definición	Ventaja principal	Limitación
Media	Promedio aritmético	Usa toda la información	Sensible a extremos
Mediana	Valor central ordenado	Robusta ante asimetría	No usa todos los valores
Moda	Valor más frecuente	Aplicable a categorías	Puede no ser única

Nota. Propiedades básicas de media, mediana y moda en datos biomédicos.

2.5.4 Elección de la medida adecuada según la distribución

La elección de la medida de tendencia central debe basarse en la forma de la distribución y en el tipo de variable. En distribuciones simétricas, como aquellas que se aproximan a la normalidad, la media y la mediana suelen coincidir o presentar valores muy cercanos. En distribuciones asimétricas, la mediana proporciona una representación más estable del centro de los datos.

En variables cualitativas, la moda suele ser la única medida aplicable. En variables ordinales, la mediana y la moda resultan más apropiadas que la media, debido a la ausencia de intervalos numéricos equivalentes entre categorías.

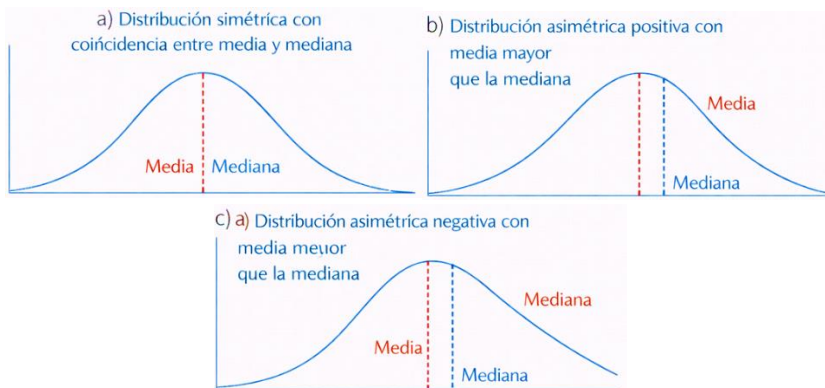


Figura 2-5. Medidas de tendencia central en diferentes distribuciones

2.5.4.1 *Ejercicio propuesto: cálculo e interpretación de medidas de tendencia central*

Planteamiento

Se registran los valores de glucosa en ayunas (mg/dL) en diez pacientes: 118, 122, 125, 130, 132, 135, 140, 145, 150, 210.

Actividad

1. Calcular la media, la mediana y la moda.
2. Interpretar los resultados desde el punto de vista clínico.

Solución

Cálculo de la media

$$\begin{aligned}\bar{x} &= \frac{118 + 122 + 125 + 130 + 132 + 135 + 140 + 145 + 150 + 210}{10} \\ &= \frac{1407}{10} = 140.7\end{aligned}$$

Cálculo de la mediana

Al ordenar los datos, los valores centrales son 132 y 135:

$$\text{Mediana} = \frac{132 + 135}{2} = 133.5$$

Cálculo de la moda

No existe moda, ya que ningún valor se repite.

Interpretación

La media es considerablemente mayor que la mediana debido al valor elevado de 210 mg/dL, que actúa como valor extremo. En esta situación, la mediana representa mejor el valor típico de la glucosa en la muestra. Desde el punto de vista clínico, esta situación ilustra la importancia de elegir la medida de tendencia central adecuada para evitar interpretaciones distorsionadas.

2.6 Medidas de dispersión: rango, varianza y desviación estándar

La descripción completa de los datos biomédicos no puede limitarse a un valor central. Dos conjuntos de datos pueden compartir la misma media o mediana y, sin embargo, diferir notablemente en su variabilidad. Las medidas de dispersión permiten cuantificar la extensión con la que los valores se distribuyen alrededor de la medida central, aportando información esencial para la interpretación clínica, la comparación entre grupos y la evaluación de la precisión de las mediciones.

En ciencias de la salud, la dispersión adquiere relevancia directa en la práctica clínica y de laboratorio. Una media de glucosa con alta variabilidad puede indicar heterogeneidad metabólica, mientras que una baja variabilidad en controles analíticos sugiere estabilidad del método. Por ello, la bioestadística

aplicada enfatiza el uso conjunto de medidas de tendencia central y de dispersión.

2.6.1 Rango

El rango es la medida de dispersión más simple y se define como la diferencia entre el valor máximo y el valor mínimo de un conjunto de datos. Proporciona una idea rápida de la amplitud total de los valores observados.

La expresión matemática del rango es:

$$R = x_{\text{máx}} - x_{\text{mín}}$$

En el ámbito clínico, el rango resulta útil como indicador preliminar de variabilidad, aunque presenta limitaciones importantes. Al depender solo de dos valores, es altamente sensible a valores extremos y no refleja la distribución interna de los datos.

Tabla 2-17. Ejemplo de cálculo del rango en valores clínicos

Variable	Valor mínimo	Valor máximo	Rango
Glucosa (mg/dL)	118	210	92

Nota. El rango resume la amplitud total de los datos, pero no describe su distribución interna.

2.6.2 Varianza

La varianza cuantifica la dispersión promedio de los valores respecto a la media. Se basa en el cálculo de las desviaciones individuales al cuadrado, lo que evita la cancelación de valores positivos y negativos. En bioestadística, se distingue entre varianza poblacional y varianza muestral, siendo esta última la más utilizada en investigación clínica.

La varianza muestral se expresa como:

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Aunque la varianza es fundamental para el desarrollo teórico de numerosos métodos estadísticos, su interpretación clínica directa es limitada debido a que se expresa en unidades cuadradas. Por esta razón, suele complementarse con la desviación estándar.

2.6.3 Desviación estándar

La desviación estándar es la raíz cuadrada de la varianza y se expresa en las mismas unidades que la variable original. Esta característica facilita su interpretación y la convierte en la medida de dispersión más utilizada en estudios biomédicos.

La desviación estándar se calcula mediante:

$$s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$$

En términos clínicos, la desviación estándar indica cuánto se alejan, en promedio, los valores individuales de la media. Valores pequeños sugieren homogeneidad, mientras que valores grandes reflejan alta variabilidad entre los sujetos o mediciones.

Tabla 2-18. Comparación de medidas de dispersión

Medida	Definición	Ventaja	Limitación
Rango	Diferencia máx–mín	Cálculo simple	Sensible a extremos
Varianza	Promedio de desviaciones cuadradas	Base analítica	Unidades cuadradas
Desviación estándar	Raíz de la varianza	Interpretación directa	Influida por extremos

Nota. Comparación de medidas de dispersión utilizadas en la descripción de datos biomédicos.

2.6.4 Interpretación clínica de la dispersión

En investigación clínica, la dispersión complementa a la media al ofrecer información sobre la heterogeneidad de la muestra. Por ejemplo, dos tratamientos pueden presentar la misma media de reducción de presión arterial, pero uno mostrar mayor variabilidad, lo que puede tener implicaciones clínicas en términos de respuesta individual.

En el laboratorio clínico, la desviación estándar se utiliza para evaluar la imprecisión analítica y forma parte del control de calidad interno. Métodos con menor dispersión ofrecen resultados más consistentes y confiables.

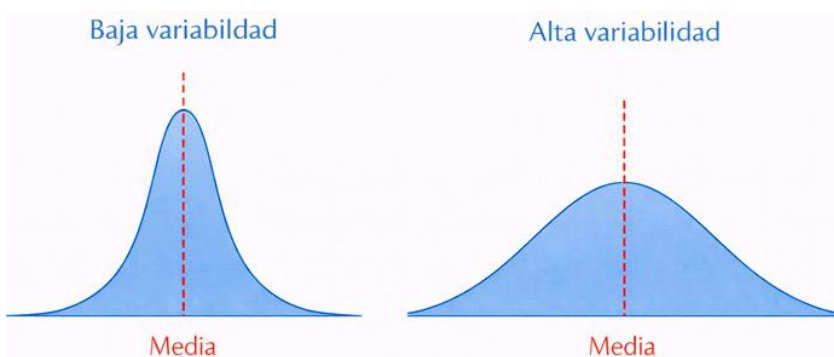


Figura 2-6. Dispersión de datos alrededor de la media

Ejercicio Propuesto: cálculo e interpretación de medidas de dispersión

Planteamiento

Se registran los valores de colesterol total (mg/dL) en ocho pacientes: 180, 185, 190, 195, 200, 205, 210, 215.

Actividad

1. Calcular el rango, la varianza y la desviación estándar.
2. Interpretar los resultados desde el punto de vista clínico.

Solución

Cálculo del rango

$$R = 215 - 180 = 35$$

Cálculo de la media

$$\bar{x} = \frac{180 + 185 + 190 + 195 + 200 + 205 + 210 + 215}{8} = \frac{1580}{8} = 197.5$$

Cálculo de la varianza

Se calculan las desviaciones respecto a la media y se elevan al cuadrado:

$$s^2 = \frac{(180 - 197.5)^2 + \dots + (215 - 197.5)^2}{7} = \frac{1050}{7} = 150$$

Cálculo de la desviación estándar

$$s = \sqrt{150} = 12.25$$

Interpretación

El rango indica una amplitud moderada de los valores de colesterol. La desviación estándar sugiere que la mayoría de los valores se sitúan aproximadamente a 12 mg/dL de la media. Desde el punto de vista clínico, esta dispersión refleja una variabilidad relativamente controlada entre los pacientes, lo que facilita la interpretación del valor promedio.

2.7 Aplicaciones en software estadístico

La organización y descripción de datos biomédicos se ha visto fortalecida de manera significativa por el uso de software estadístico, el cual permite procesar grandes volúmenes de información con precisión, reproducibilidad y eficiencia. En la formación y práctica profesional de las ciencias de la salud, el uso de herramientas computacionales no sustituye la comprensión conceptual de la bioestadística, sino que la complementa, facilitando la aplicación rigurosa de métodos descriptivos y la correcta interpretación clínica de los resultados.

Entre los entornos más utilizados en investigación biomédica se encuentran SPSS, R y Python. Cada uno presenta características particulares, pero todos permiten realizar tabulación, construcción de tablas de frecuencia, representación gráfica y cálculo de medidas descriptivas de forma sistemática y verificable.

2.7.1 **Uso de software estadístico en la descripción de datos clínicos**

El empleo de software estadístico en salud responde a varias necesidades fundamentales:

- Reducción de errores de cálculo manual.
- Manejo eficiente de bases de datos extensas.
- Reproducibilidad de los análisis.
- Generación automática de tablas y gráficos estandarizados.

La metodología en investigación clínica puntualiza que el uso de software debe ir acompañado de una adecuada comprensión de los supuestos estadísticos, ya que una mala interpretación de los resultados computacionales puede generar conclusiones incorrectas, aun cuando los cálculos sean técnicamente correctos.

2.7.2 **Aplicaciones descriptivas en SPSS**

SPSS es un software ampliamente utilizado en el ámbito de las ciencias de la salud, especialmente en estudios clínicos, epidemiológicos y educativos. Su entorno gráfico facilita la tabulación y el análisis descriptivo de datos clínicos sin requerir conocimientos avanzados de programación.

Con SPSS es posible:

- Crear bases de datos estructuradas con variables clínicas y de laboratorio.
- Generar tablas de frecuencia para variables cualitativas y cuantitativas.
- Calcular media, mediana, moda, rango y desviación estándar.
- Construir diagramas de barras, histogramas y gráficos de dispersión.

Tabla 2-19. Funciones descriptivas habituales en SPSS

Función	Resultado
Frecuencias	Tablas y porcentajes
Descriptivos	Media y dispersión
Gráficos	Barras, histogramas

Nota. Principales opciones de SPSS para análisis descriptivo de datos biomédicos.

La fortaleza de SPSS radica en su facilidad de uso y en la estandarización de salidas, lo que lo hace adecuado para prácticas clínicas y formativas.

2.7.3 Aplicaciones descriptivas en R

R es un entorno estadístico basado en programación que ofrece gran flexibilidad y potencia analítica. Es ampliamente utilizado en investigación biomédica avanzada y permite una descripción detallada y reproducible de los datos (Dalgaard, 2008).

En R, las operaciones descriptivas se realizan mediante funciones específicas, lo que exige una comprensión más profunda de la estructura de los datos. Sin embargo, esta característica favorece la transparencia del análisis y la trazabilidad de cada procedimiento estadístico.

Ejemplos de funciones descriptivas en R:

`mean(x),median(x),sd(x)`

Estas funciones permiten calcular de forma directa las medidas de tendencia central y dispersión previamente estudiadas.

Tabla 2-20. Funciones básicas de R para estadística descriptiva

Función	Descripción
<code>mean()</code>	Media aritmética
<code>median()</code>	Mediana
<code>sd()</code>	Desviación estándar
<code>summary()</code>	Resumen estadístico

Nota. Funciones fundamentales de R utilizadas en la descripción inicial de datos biomédicos.

El uso de R favorece la reproducibilidad del análisis, aspecto altamente valorado en la investigación científica contemporánea.

2.7.4 Aplicaciones descriptivas en Python

Python se ha consolidado como una herramienta versátil para el análisis de datos en ciencias de la salud, integrando bioestadística, análisis exploratorio y visualización. Mediante bibliotecas especializadas, Python permite realizar análisis descriptivos con gran claridad y eficiencia (Mckinney, 2022).

Las bibliotecas más empleadas para estadística descriptiva incluyen estructuras de datos tabulares y funciones matemáticas que facilitan el manejo de información clínica y de laboratorio.

Ejemplo conceptual de cálculo descriptivo:

$$\text{media} = \frac{1}{n} \sum x_i$$

El software ejecuta este cálculo de manera automática, minimizando errores aritméticos y permitiendo trabajar con bases de datos extensas.

Tabla 2-21. Operaciones descriptivas en Python

Operación	Resultado
describe	Resumen estadístico
mean	Media
std	Desviación estándar
plot	Gráficos básicos

Nota. Operaciones comunes en Python para describir datos biomédicos de forma eficiente.

Python destaca por su capacidad de integración con técnicas más avanzadas, como análisis predictivo y aprendizaje automático, que se abordarán en capítulos posteriores.

2.7.5 Comparación didáctica entre SPSS, R y Python

Cada software presenta ventajas particulares según la aplicación práctica. En la enseñanza de las ciencias de la salud, la selección debe considerar los objetivos formativos, el nivel del estudiante y el tipo de análisis requerido.

Tabla 2-22. Comparación general de software estadístico en salud

Software	Nivel de programación	Uso frecuente
SPSS	Bajo	Estudios clínicos descriptivos
R	Medio	Investigación biomédica
Python	Medio	Análisis integrado y automatización

Nota. Comparación general de herramientas estadísticas según nivel de complejidad y aplicación en salud.

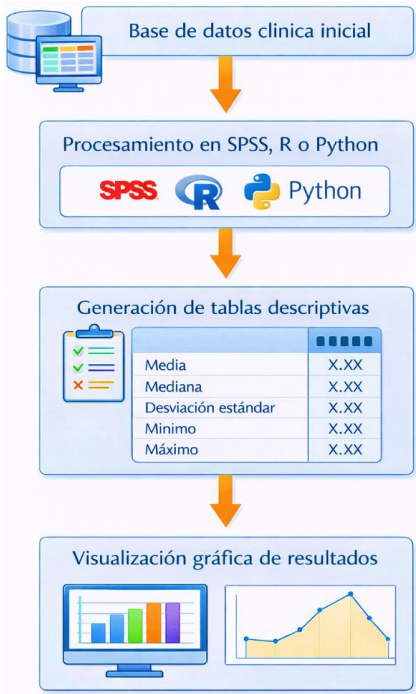


Figura 2-7. Flujo de análisis descriptivo con software estadístico

Ejercicio aplicado: descripción de datos con apoyo de software

Planteamiento

Se dispone de una base de datos con valores de presión arterial sistólica (mmHg) de diez pacientes:

118, 122, 130, 135, 138, 140, 145, 150, 155, 160.

Actividad

1. Calcular media, mediana y desviación estándar utilizando un software estadístico.
2. Elaborar un histograma de los datos.
3. Interpretar los resultados desde el punto de vista clínico.

Solución:

El software calcula automáticamente:

- Media aproximada de 139.3 mmHg.
- Mediana de 139 mmHg.
- Desviación estándar moderada que indica variabilidad clínica relevante.

El histograma muestra concentración de valores en rangos elevados, lo que sugiere predominio de cifras compatibles con hipertensión.

Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

CAPÍTULO III

Probabilidad y Distribuciones en Salud

AUTOR: Jijón Cañarte Lidia Fernanda



3 Probabilidad y Distribuciones en Salud

3.1 Concepto de probabilidad en el contexto sanitario

La probabilidad constituye un pilar fundamental de la bioestadística y de la toma de decisiones en ciencias de la salud. En términos generales, la probabilidad expresa el grado de incertidumbre asociado a la ocurrencia de un evento. En el ámbito sanitario, esta incertidumbre está presente en el diagnóstico, el pronóstico, la respuesta a tratamientos, la aparición de eventos adversos y la evolución de enfermedades. La bioestadística utiliza el lenguaje probabilístico para cuantificar esa incertidumbre y transformarla en información útil para la práctica clínica, la investigación biomédica y la salud pública (Rosner, 2016; Gordis, 2014).

Desde una perspectiva aplicada, la probabilidad no se limita a un concepto matemático abstracto. En medicina y disciplinas afines, se interpreta como una medida cuantitativa del riesgo o de la verosimilitud de un evento sanitario, como la probabilidad de que un paciente presente una enfermedad, de que una prueba diagnóstica resulte positiva o de que una intervención produzca un determinado desenlace clínico (Altman, 1991; Motulsky, 2015).

3.1.1 Definición y enfoques de la probabilidad

En bioestadística aplicada a la salud, la probabilidad se define como un valor numérico comprendido entre 0 y 1 que indica la posibilidad de ocurrencia de un evento. Un valor cercano a 0 indica baja probabilidad, mientras que un valor cercano a 1 indica alta probabilidad. Este concepto permite comparar riesgos y fundamentar decisiones clínicas basadas en evidencia cuantitativa (W. Daniel & Cross, 2013).

Existen distintos enfoques conceptuales para interpretar la probabilidad. En la práctica sanitaria, los más utilizados son el enfoque clásico y el enfoque frecuentista. El enfoque clásico se basa en la relación entre casos favorables y casos posibles cuando todos los resultados son igualmente probables. El enfoque frecuentista, ampliamente adoptado en epidemiología y medicina, define la probabilidad como la frecuencia relativa de un evento cuando el experimento se repite un gran número de veces.

Este último enfoque resulta especialmente pertinente en salud, ya que muchas probabilidades se estiman a partir de datos observados en poblaciones, como tasas de incidencia, prevalencia o proporciones de resultados clínicos.

3.1.2 Probabilidad y eventos sanitarios

Un evento sanitario es cualquier resultado observable de interés clínico o epidemiológico. Ejemplos comunes incluyen la presencia de una enfermedad, la ocurrencia de una complicación, la respuesta favorable a un tratamiento o la obtención de un resultado positivo en una prueba diagnóstica. La probabilidad de un evento se calcula como la proporción de veces que dicho evento ocurre en relación con el total de observaciones realizadas.

La expresión básica de la probabilidad de un evento A es:

$$P(A) = \frac{\text{Número de casos favorables}}{\text{Número total de casos}}$$

En investigación clínica, esta expresión se aplica para estimar riesgos, proporciones y tasas, siempre considerando el tamaño de la muestra y la población del estudio.

Tabla 3-1. Ejemplos de eventos sanitarios y su interpretación probabilística

Evento sanitario	Interpretación de la probabilidad
Prueba diagnóstica positiva	Probabilidad de detectar la enfermedad
Evento adverso	Probabilidad de efecto no deseado
Respuesta al tratamiento	Probabilidad de mejoría clínica

Nota. Ejemplos de eventos sanitarios donde la probabilidad cuantifica la ocurrencia de resultados clínicamente relevantes.

3.1.3 Probabilidad y riesgo en salud

En el lenguaje sanitario, el término riesgo se utiliza con frecuencia como sinónimo de probabilidad, aunque conceptualmente se refiere a la probabilidad de que ocurra un evento adverso en un periodo determinado. La cuantificación

del riesgo permite comparar poblaciones, evaluar intervenciones preventivas y estimar beneficios y daños potenciales.

Por ejemplo, si en una cohorte de 1 000 personas 50 desarrollan una enfermedad durante un año, el riesgo anual se estima como 0.05. Esta medida probabilística se convierte en un insumo para la planificación sanitaria y la evaluación de estrategias de control.

3.1.4 Interpretación clínica de la probabilidad

La correcta interpretación de la probabilidad en salud requiere considerar la práctica clínica y la población estudiada. Una probabilidad elevada no implica certeza absoluta, ni una probabilidad baja descarta completamente la ocurrencia de un evento. La probabilidad expresa incertidumbre cuantificada y debe comunicarse de manera clara a pacientes y profesionales de la salud.

En la práctica clínica, la probabilidad se integra con el juicio profesional, la experiencia clínica y las preferencias del paciente. Esta integración constituye la base de la medicina basada en evidencia, donde las decisiones se apoyan en estimaciones probabilísticas derivadas de estudios bien diseñados.

Tabla 3-2. Interpretación cualitativa de valores de probabilidad en salud

Probabilidad	Interpretación clínica aproximada
< 0.10	Evento poco frecuente
0.10 – 0.50	Evento posible
> 0.50	Evento probable

Nota. Guía orientativa para la interpretación clínica de valores de probabilidad en decisiones sanitarias.

3.1.5 Probabilidad y toma de decisiones sanitarias

La probabilidad desempeña un rol central en la toma de decisiones diagnósticas y terapéuticas. Por ejemplo, la decisión de solicitar una prueba complementaria depende de la probabilidad previa de enfermedad, mientras que la indicación de un tratamiento puede variar según la probabilidad estimada de beneficio frente al riesgo de efectos adversos.

Desde el punto de vista bioestadístico, esta probabilidad previa se actualiza a medida que se dispone de nueva información clínica o diagnóstica, proceso que será desarrollado con mayor profundidad en el epígrafe dedicado al teorema de Bayes.



Figura 3-1. Probabilidad y decisión clínica

Ejercicio: cálculo e interpretación de probabilidad en un ambiente sanitario

Planteamiento

En un centro de salud se evaluaron 200 pacientes con sospecha de infección respiratoria. De ellos, 40 fueron confirmados mediante pruebas de laboratorio.

Actividad

1. Calcular la probabilidad de confirmación de infección en esta población.
2. Interpretar el resultado desde el punto de vista clínico.

Solución

Cálculo de la probabilidad

$$P(\text{Infección}) = \frac{40}{200} = 0.20$$

Interpretación

La probabilidad estimada de infección en la población evaluada es de 0.20, lo que indica que dos de cada diez pacientes presentan la enfermedad. Desde el punto de vista clínico, este valor sugiere una probabilidad moderada, que justifica la realización de pruebas diagnósticas confirmatorias y la aplicación de medidas preventivas en el grupo evaluado.

3.2 Regla de adición y multiplicación de probabilidades

Las reglas de adición y multiplicación constituyen principios operativos básicos de la probabilidad y resultan indispensables para el análisis de eventos sanitarios que pueden ocurrir de forma conjunta o alternativa. En ciencias de la salud, estas reglas se aplican con frecuencia para evaluar combinaciones de diagnósticos, resultados de pruebas, factores de riesgo concurrentes y desenlaces clínicos múltiples. Su correcta comprensión permite estructurar razonamientos probabilísticos coherentes y evitar interpretaciones erróneas en la práctica clínica y en la investigación biomédica.

Desde una perspectiva aplicada, estas reglas facilitan el cálculo de probabilidades en situaciones reales donde los eventos no ocurren de manera aislada, sino que se relacionan entre sí de diversas formas. La bioestadística traduce estas relaciones en expresiones matemáticas simples, pero conceptualmente potentes, que sustentan decisiones diagnósticas y epidemiológicas fundamentadas (Altman, 1991).

3.2.1 Regla de adición de probabilidades

La regla de adición se utiliza cuando se desea calcular la probabilidad de que ocurra al menos uno de dos eventos. En la práctica sanitaria, esto puede referirse a la probabilidad de que un paciente presente una de varias enfermedades, o a que una prueba diagnóstica arroje uno u otro resultado de interés (Gordis, 2014).

Si los eventos A y B son mutuamente excluyentes, es decir, no pueden ocurrir simultáneamente, la probabilidad de que ocurra A o B se calcula como:

$$P(A \cup B) = P(A) + P(B)$$

Por ejemplo, si se estudia la probabilidad de que un paciente tenga enfermedad A o enfermedad B, y se sabe que no puede presentar ambas al mismo tiempo, la probabilidad total se obtiene sumando las probabilidades individuales.

3.2.2 Regla de adición para eventos no excluyentes

En muchas situaciones clínicas, los eventos no son mutuamente excluyentes. Un paciente puede presentar dos condiciones simultáneamente, como hipertensión y diabetes. En estos casos, aplicar la suma directa conduciría a una sobreestimación de la probabilidad, ya que se contaría dos veces la ocurrencia conjunta (Rothman et al., 2008).

La regla general de adición para eventos no excluyentes es:

$$P(A \cup B) = P(A) + P(B) - P(A \cap B)$$

donde $P(A \cap B)$ representa la probabilidad de que ambos eventos ocurran simultáneamente.

Tabla 3-3. Regla de adición de probabilidades en eventos sanitarios

Tipo de eventos	Expresión matemática	Ejemplo clínico
Mutuamente excluyentes	$P(A)+P(B)$	Diagnóstico A o B
No excluyentes	$P(A)+P(B)-P(A\cap B)$	Comorbilidades

Nota. Aplicación de la regla de adición según la relación entre eventos clínicos o epidemiológicos.

3.2.3 Regla de multiplicación de probabilidades

La regla de multiplicación se emplea para calcular la probabilidad de que dos eventos ocurran de manera conjunta. En el ámbito sanitario, esta regla resulta esencial para evaluar secuencias diagnósticas, procesos terapéuticos y combinaciones de factores de riesgo (Pagano y Gauvreau, 2018).

Cuando los eventos A y B son independientes, es decir, la ocurrencia de uno no afecta la probabilidad del otro, la probabilidad conjunta se calcula como:

$$P(A \cap B) = P(A) \times P(B)$$

Un ejemplo clínico sería la probabilidad de que un paciente seleccionado al azar sea hombre y presente un resultado normal en una prueba, siempre que ambas condiciones sean independientes.

3.2.4 Multiplicación en eventos dependientes

En la práctica clínica, muchos eventos son dependientes. Por ejemplo, la probabilidad de un resultado positivo en una prueba confirmatoria depende del resultado de una prueba previa. En estos casos, se utiliza la probabilidad condicional (Daniel y Cross, 2017).

La expresión general es:

$$P(A \cap B) = P(A) \times P(B | A)$$

donde $P(B | A)$ representa la probabilidad de que ocurra B dado que ya ocurrió A.

Este enfoque es especialmente relevante en procesos diagnósticos escalonados y constituye la base conceptual para el teorema de Bayes, que se desarrollará en el epígrafe siguiente.

Tabla 3-4. Regla de multiplicación

Tipo de eventos	Expresión	Aplicación
Independientes	$P(A) \times P(B)$	Eventos no relacionados
Dependientes	$P(A) \times P(B A)$	

Nota. Uso de la regla de multiplicación según la dependencia entre eventos clínicos o diagnósticos.

3.2.5 Interpretación sanitaria de las reglas de probabilidad

La correcta aplicación de las reglas de adición y multiplicación permite estimar probabilidades complejas a partir de información básica. En epidemiología, estas reglas se emplean para analizar la coexistencia de factores de riesgo, mientras que en diagnóstico clínico se utilizan para evaluar rutas diagnósticas y combinaciones de pruebas.

Una aplicación incorrecta de estas reglas puede conducir a conclusiones erróneas, como sobreestimar riesgos o subestimar la probabilidad de eventos concurrentes. Por ello, la bioestadística aplicada insiste en identificar primero la relación entre los eventos antes de seleccionar la expresión matemática correspondiente.

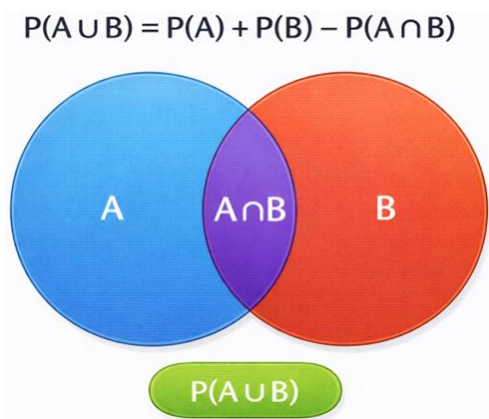


Figura 3-2. Representación gráfica de la regla de adición

Ejercicio: reglas de adición y multiplicación en salud

Planteamiento

En una población adulta se observa lo siguiente:

- Probabilidad de hipertensión: 0.30
- Probabilidad de diabetes: 0.20
- Probabilidad de presentar ambas condiciones: 0.10

Actividad

1. Calcular la probabilidad de que un individuo presente hipertensión o diabetes.
2. Interpretar el resultado desde el punto de vista sanitario.

Solución

Cálculo mediante la regla de adición

$$P(H \cup D) = 0.30 + 0.20 - 0.10 = 0.40$$

Interpretación

La probabilidad de que un individuo presente al menos una de las dos condiciones es de 0.40. Desde el punto de vista sanitario, este resultado indica que cuatro de cada diez adultos presentan hipertensión, diabetes o ambas, lo que evidencia una carga relevante de enfermedad crónica en la población estudiada.

3.3 Teorema de Bayes aplicado al diagnóstico clínico

El teorema de Bayes constituye uno de los fundamentos matemáticos más relevantes para el razonamiento diagnóstico en ciencias de la salud. Su aporte principal radica en permitir la actualización de la probabilidad de una enfermedad a partir de nueva información, como el resultado de una prueba diagnóstica. En la práctica clínica, los profesionales no toman decisiones basándose únicamente en resultados aislados, sino integrando la probabilidad previa de enfermedad con la evidencia aportada por pruebas complementarias. El teorema de Bayes formaliza este proceso de actualización probabilística y lo traduce en un marco cuantitativo riguroso.

La epidemiología clínica ha destacado que muchos errores diagnósticos se originan en la interpretación inadecuada de pruebas diagnósticas, especialmente cuando no se considera la probabilidad previa de enfermedad. El enfoque bayesiano permite evitar estas interpretaciones erróneas al situar los resultados diagnósticos dentro de un marco probabilístico coherente.



Figura 3-3. Actualización bayesiana en el proceso diagnóstico

3.3.1 Formulación general del teorema de Bayes

Desde el punto de vista matemático, el teorema de Bayes relaciona la probabilidad condicional de un evento A dado que ha ocurrido un evento B, con probabilidades inversas conocidas. Su expresión general es:

$$P(A | B) = \frac{P(B | A) \cdot P(A)}{P(B)}$$

En el entorno sanitario:

- A representa la presencia de una enfermedad.
- B representa un resultado positivo en una prueba diagnóstica.

Así, $P(A | B)$ corresponde a la probabilidad de que un paciente tenga la enfermedad dado que la prueba resultó positiva, magnitud conocida como probabilidad posterior.

3.3.2 Componentes bayesianos en el diagnóstico clínico

La aplicación del teorema de Bayes al diagnóstico clínico se apoya en cuatro componentes fundamentales:

- Probabilidad previa de enfermedad, estimada a partir de la prevalencia o del juicio clínico.
- Sensibilidad de la prueba, que corresponde a la probabilidad de obtener un resultado positivo en personas con la enfermedad.
- Especificidad de la prueba, que corresponde a la probabilidad de obtener un resultado negativo en personas sin la enfermedad.
- Probabilidad posterior, que resulta de combinar la información previa con el resultado de la prueba.

Estos elementos permiten comprender que el valor diagnóstico de una prueba no depende solo de sus propiedades técnicas, sino también de la frecuencia de la enfermedad en la población evaluada.

Tabla 3-5. Componentes del teorema de Bayes en diagnóstico clínico

Componente	Definición	Interpretación clínica
Probabilidad previa	Probabilidad inicial de enfermedad	Sospecha clínica
Sensibilidad	$P(\text{Prueba} +)$	Enfermedad)
Especificidad	$P(\text{Prueba} -)$	No enfermedad)
Probabilidad posterior	$P(\text{Enfermedad})$	Prueba +)

Nota. Elementos que intervienen en la actualización probabilística del diagnóstico clínico mediante el teorema de Bayes.

3.3.3 Interpretación clínica del teorema de Bayes

En la práctica clínica, el teorema de Bayes explica por qué una prueba con alta sensibilidad y especificidad puede tener un bajo valor confirmatorio cuando se aplica en poblaciones con baja prevalencia de la enfermedad. Este fenómeno es especialmente relevante en programas de tamizaje y en el uso indiscriminado de pruebas diagnósticas.

Por ejemplo, una prueba aplicada a una población con baja probabilidad previa puede generar un número considerable de resultados positivos que no corresponden a casos reales, lo que incrementa la carga diagnóstica y los costos sanitarios. La bioestadística aplicada enfatiza que el razonamiento bayesiano debe formar parte de la formación clínica para mejorar la toma de decisiones diagnósticas.

3.3.4 Representación tabular del razonamiento bayesiano

Una forma didáctica y ampliamente utilizada para aplicar el teorema de Bayes en salud es mediante tablas de contingencia. Estas tablas permiten visualizar la relación entre el estado real de la enfermedad y los resultados de la prueba diagnóstica (Pagano y Gauvreau, 2018).

Tabla 3-6. Tabla de contingencia para una prueba diagnóstica

Estado clínico	Prueba positiva	Prueba negativa	Total
Enfermedad presente	Verdadero positivo	Falso negativo	Total enfermos
Enfermedad ausente	Falso positivo	Verdadero negativo	Total sanos
Total	Total positivos	Total negativos	Total

Nota. Estructura básica para visualizar resultados diagnósticos y aplicar el razonamiento bayesiano.

3.3.5 Aplicación práctica del teorema de Bayes

Supóngase una enfermedad con una prevalencia del 5 por ciento en una población. Se dispone de una prueba diagnóstica con sensibilidad del 90 por

ciento y especificidad del 95 por ciento. El interés clínico radica en conocer la probabilidad real de enfermedad cuando la prueba resulta positiva.

Cálculo paso a paso:

- Probabilidad previa de enfermedad:

$$P(E) = 0.05$$

Probabilidad de no enfermedad:

$$P(\neg E) = 0.95$$

- Probabilidad de resultado positivo:

$$P(+) = P(+|E) \cdot P(E) + P(+|\neg E) \cdot P(\neg E)$$

$$P(+) = 0.90 \times 0.05 + 0.05 \times 0.95 = 0.045 + 0.0475 = 0.0925$$

- Probabilidad posterior de enfermedad dado resultado positivo:

$$P(E|+) = \frac{0.90 \times 0.05}{0.0925} = 0.486$$

Interpretación clínica

Aunque la prueba presenta alta sensibilidad y especificidad, la probabilidad de que un paciente con resultado positivo tenga realmente la enfermedad es de aproximadamente 0.49. Este resultado evidencia la importancia de considerar la probabilidad previa y evita interpretaciones diagnósticas excesivamente confiadas basadas solo en el resultado de la prueba.

Ejercicio: razonamiento bayesiano en diagnóstico clínico

Planteamiento

Una enfermedad tiene una prevalencia del 10 por ciento en una población. Una prueba diagnóstica presenta sensibilidad del 85 por ciento y especificidad del 90 por ciento.

Actividad

1. Calcular la probabilidad de que un paciente tenga la enfermedad si la prueba resulta positiva.

2. Interpretar el resultado desde el punto de vista clínico.

Solución

$$P(E) = 0.10, P(\neg E) = 0.90$$

$$P(+) = 0.85 \times 0.10 + 0.10 \times 0.90 = 0.085 + 0.09 = 0.175$$

$$P(E | +) = \frac{0.85 \times 0.10}{0.175} = 0.486$$

Interpretación

La probabilidad posterior de enfermedad es cercana al 49 por ciento. Desde el punto de vista clínico, este resultado indica que un resultado positivo no es concluyente por sí solo y debe interpretarse junto con la evaluación clínica y pruebas adicionales.

3.4 Distribución binomial y aplicación a resultados positivos y negativos

La distribución binomial es un modelo probabilístico fundamental para describir fenómenos sanitarios en los que cada observación puede clasificarse en uno de dos resultados mutuamente excluyentes, como positivo o negativo, éxito o fracaso, presencia o ausencia de enfermedad. En ciencias de la salud, esta distribución se utiliza ampliamente para analizar resultados de pruebas diagnósticas, respuestas a tratamientos, eventos adversos y desenlaces clínicos dicotómicos. Su aplicación permite estimar probabilidades, evaluar incertidumbre y fundamentar decisiones clínicas y epidemiológicas basadas en datos.

Desde una perspectiva aplicada, la distribución binomial conecta el razonamiento probabilístico con situaciones clínicas reales, donde se observa un número finito de ensayos bajo condiciones similares. Este enfoque resulta especialmente pertinente en estudios de diagnóstico, ensayos clínicos y evaluaciones de calidad en el laboratorio clínico.

3.4.1 Condiciones de aplicación de la distribución binomial

Para que un fenómeno sanitario pueda modelarse mediante una distribución binomial, deben cumplirse las siguientes condiciones:

1. El experimento consta de un número fijo de ensayos n .
2. Cada ensayo tiene dos resultados posibles.
3. La probabilidad de éxito p es constante en todos los ensayos.
4. Los ensayos son independientes entre sí.

En la práctica sanitaria, un ensayo puede interpretarse como la aplicación de una prueba diagnóstica a un paciente o la observación de un desenlace clínico en un periodo definido.

3.4.2 Función de probabilidad binomial

La probabilidad de obtener exactamente k éxitos en n ensayos se calcula mediante la función de probabilidad binomial:

$$P(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}$$

donde:

- X es la variable aleatoria que representa el número de éxitos,
- n es el número total de ensayos,
- k es el número de éxitos observados,
- p es la probabilidad de éxito en cada ensayo.

Este modelo permite cuantificar la probabilidad de observar un número determinado de resultados positivos en un conjunto de pruebas clínicas o de laboratorio.

3.4.3 Interpretación sanitaria de los parámetros binomiales

En aplicaciones clínicas, el parámetro p suele representar una proporción de interés, como la prevalencia de una enfermedad, la tasa de respuesta a un tratamiento o la probabilidad de obtener un resultado positivo en una prueba. El número de ensayos n corresponde al tamaño de la muestra o al número de pacientes evaluados.

La media y la varianza de la distribución binomial se expresan como:

$$\mu = np$$
$$\sigma^2 = np(1 - p)$$

Estas expresiones permiten anticipar el comportamiento esperado de los resultados y evaluar la variabilidad asociada al proceso clínico o diagnóstico.

Tabla 3-7. Parámetros e interpretación clínica de la distribución binomial

Parámetro	Significado	Interpretación sanitaria
n	Número de ensayos	Pacientes evaluados
p	Probabilidad de éxito	Positividad esperada
k	Éxitos observados	Casos positivos
$1 - p$	Probabilidad de fracaso	Casos negativos

Nota. Correspondencia entre los parámetros de la distribución binomial y situaciones clínicas habituales.

3.4.4 Aplicación a resultados de pruebas diagnósticas

La distribución binomial resulta especialmente útil para modelar resultados positivos y negativos de pruebas diagnósticas en un grupo de pacientes. Por ejemplo, si se conoce la probabilidad de positividad de una prueba y se aplica a un número fijo de individuos, la distribución binomial permite estimar la probabilidad de observar cierto número de resultados positivos.

Este enfoque se utiliza en el diseño de estudios diagnósticos, en la evaluación de programas de tamizaje y en el control de calidad de pruebas de laboratorio, donde se espera una proporción determinada de resultados positivos bajo condiciones controladas.

Tabla 3-8. Ejemplo de resultados binomiales en pruebas diagnósticas

Número de pruebas	Positivos esperados	Negativos esperados
10	$10p$	$10(1 - p)$
50	$50p$	$50(1 - p)$

Nota. Expectativa de resultados positivos y negativos según la probabilidad de éxito en pruebas diagnósticas.

3.4.5 Aproximación clínica y límites del modelo binomial

Aunque la distribución binomial es conceptualmente clara, su aplicación práctica requiere verificar que se cumplan las condiciones de independencia y probabilidad constante. En la práctica clínica, estas condiciones pueden verse afectadas por factores como heterogeneidad de pacientes, cambios en la técnica diagnóstica o dependencia temporal entre observaciones.

Cuando el tamaño muestral es grande y la probabilidad de éxito es moderada, la distribución binomial puede aproximarse mediante la distribución normal, aspecto que se desarrollará en el siguiente epígrafe al analizar su relevancia en valores de referencia de laboratorio.



Figura 3-4. Distribución binomial de resultados positivos y negativos

Ejercicio: distribución binomial en diagnóstico clínico

Planteamiento

Una prueba diagnóstica tiene una probabilidad de resultado positivo del 20 por ciento. Se aplica la prueba a 10 pacientes independientes.

Actividad

1. Calcular la probabilidad de obtener exactamente 3 resultados positivos.
2. Determinar el número esperado de resultados positivos.
3. Interpretar los resultados desde el punto de vista sanitario.

Solución

Cálculo de la probabilidad

$$P(X = 3) = \binom{10}{3} (0.20)^3 (0.80)^7$$
$$P(X = 3) = 120 \times 0.008 \times 0.2097 = 0.201$$

Número esperado de resultados positivos

$$\mu = np = 10 \times 0.20 = 2$$

Interpretación

La probabilidad de obtener exactamente tres resultados positivos es cercana a 0.20, lo que indica que este resultado es compatible con el comportamiento esperado de la prueba. El número promedio esperado de positivos es dos, por lo que observar tres resultados positivos no representa una situación inusual desde el punto de vista clínico.

3.5 Distribución normal y su relevancia en valores de referencia de laboratorio

La distribución normal ocupa un lugar central en la bioestadística aplicada a las ciencias de la salud debido a su capacidad para modelar, de manera aproximada, la variabilidad natural de numerosas variables biológicas y de laboratorio. Mediciones como concentraciones bioquímicas, parámetros hematológicos, presión arterial y características antropométricas suelen presentar distribuciones simétricas alrededor de un valor central, lo que justifica el uso de este modelo probabilístico en la práctica clínica y diagnóstica.

Desde el punto de vista sanitario, la relevancia de la distribución normal radica en su aplicación directa a la definición de valores de referencia de laboratorio, al control de calidad analítico y a la interpretación clínica de resultados individuales en relación con una población de referencia.

3.5.1 Características fundamentales de la distribución normal

La distribución normal se caracteriza por una forma simétrica en torno a su media, con una curva continua conocida como campana de Gauss. Esta distribución queda completamente definida por dos parámetros: la media μ , que indica el centro de la distribución, y la desviación estándar σ , que mide la dispersión de los valores alrededor de la media.

Entre sus propiedades más relevantes se destacan:

- Simetría respecto a la media.
- Coincidencia entre media, mediana y moda.
- Mayor concentración de valores alrededor del centro.

Estas propiedades permiten utilizar la media como un resumen adecuado del valor típico cuando la distribución es aproximadamente normal.

3.5.2 Regla empírica y su interpretación clínica

Una propiedad ampliamente utilizada de la distribución normal es la regla empírica, que describe la proporción de datos que se concentran dentro de determinados intervalos alrededor de la media:

- Aproximadamente el 68 por ciento de los valores se encuentran dentro de $\mu \pm 1\sigma$.
- Aproximadamente el 95 por ciento de los valores se encuentran dentro de $\mu \pm 2\sigma$.
- Aproximadamente el 99.7 por ciento de los valores se encuentran dentro de $\mu \pm 3\sigma$.

En el laboratorio clínico, esta regla se emplea para evaluar si un resultado individual se encuentra dentro del rango esperado para una población sana o si se aparta de manera significativa (Altman, 1991).

Tabla 3-9. Regla empírica de la distribución normal

Intervalo alrededor de la media	Porcentaje aproximado
$\mu \pm 1\sigma$	68 %
$\mu \pm 2\sigma$	95 %
$\mu \pm 3\sigma$	99.7 %

Nota. Proporción aproximada de valores contenidos en intervalos simétricos alrededor de la media en una distribución normal.

3.5.3 Distribución normal y valores de referencia de laboratorio

Los valores de referencia de laboratorio se definen como los intervalos que abarcan los resultados esperados en una población aparentemente sana. Tradicionalmente, estos intervalos se establecen considerando el 95 por ciento central de la distribución, lo que corresponde aproximadamente a $\mu \pm 2\sigma$ cuando los datos siguen una distribución normal.

Por ejemplo, si la concentración media de un analito es de 100 unidades con una desviación estándar de 10 unidades, el intervalo de referencia aproximado se situaría entre 80 y 120 unidades. Este enfoque permite clasificar resultados como normales o fuera de rango de manera objetiva y reproducible.

3.5.4 Estandarización y puntuación z

La distribución normal también permite estandarizar valores individuales mediante la puntuación z, que expresa cuántas desviaciones estándar se encuentra un valor respecto a la media. La puntuación z se calcula como:

$$z = \frac{x - \mu}{\sigma}$$

Esta transformación facilita la comparación de resultados obtenidos en diferentes escalas o situaciones, aspecto relevante en estudios multicéntricos y en la evaluación longitudinal de pacientes.

Tabla 3-10. Interpretación clínica de la puntuación z

Valor de z	Interpretación aproximada
0	Igual a la media
± 1	Variación esperada
± 2	Valor límite
± 3	Valor inusual

Nota. Interpretación clínica general de la posición de un valor individual dentro de una distribución normal.

3.5.5 Limitaciones del uso de la distribución normal

Aunque la distribución normal es ampliamente utilizada, no todas las variables biomédicas siguen este patrón. Variables como tiempos de hospitalización, concentraciones hormonales o marcadores inflamatorios suelen presentar asimetría. En estos casos, el uso directo de la media y la desviación estándar puede resultar inapropiado.

Por ello, la bioestadística aplicada recomienda verificar la forma de la distribución mediante gráficos y medidas descriptivas antes de asumir normalidad, y considerar transformaciones o métodos alternativos cuando la suposición no se cumple.

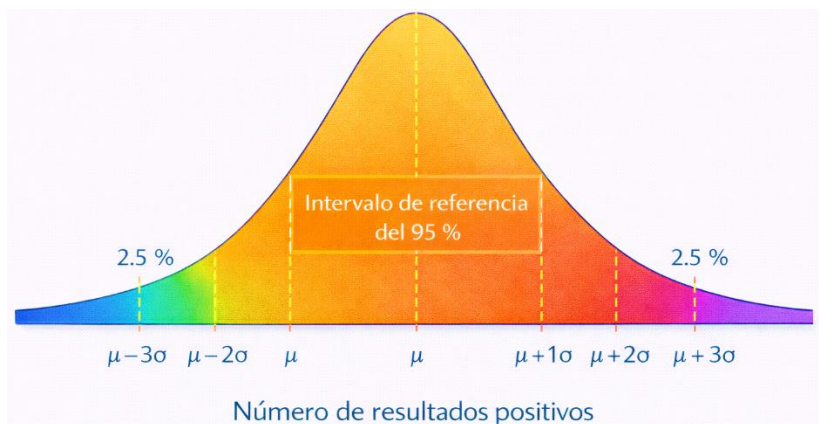


Figura 3-5. Distribución normal y valores de referencia

Propuesta de figura con una curva normal que muestre la media en el centro y los intervalos $\mu \pm 1\sigma$, $\mu \pm 2\sigma$ y $\mu \pm 3\sigma$, destacando el intervalo de referencia del 95 por ciento.

Ejercicio: valores de referencia basados en la distribución normal

Planteamiento

En una población sana, la concentración de sodio sérico presenta una media de 140 mmol/L y una desviación estándar de 4 mmol/L.

Actividad

1. Calcular el intervalo de referencia aproximado del 95 por ciento.
2. Determinar la puntuación z para un paciente con sodio sérico de 148 mmol/L.
3. Interpretar los resultados clínicamente.

Solución

Intervalo de referencia

$$\mu \pm 2\sigma = 140 \pm 2(4) = [132, 148]$$

Cálculo de la puntuación z

$$z = \frac{148 - 140}{4} = 2$$

Interpretación

El valor de sodio de 148 mmol/L se sitúa en el límite superior del intervalo de referencia. Desde el punto de vista clínico, este resultado requiere vigilancia y correlación con la situación del paciente, aunque aún se encuentra dentro del rango esperado para la población de referencia.

3.6 Distribuciones t de Student, Chi-cuadrado y F de Fisher

En la investigación biomédica y clínica, no siempre es posible aplicar la distribución normal de forma directa. En muchas situaciones, el tamaño muestral es reducido, la varianza poblacional es desconocida o se requiere analizar relaciones entre variables categóricas o comparar variabilidades entre grupos. Para estos escenarios, la bioestadística dispone de otras distribuciones fundamentales que amplían el marco probabilístico aplicado a la salud. Entre las más utilizadas se encuentran la distribución t de Student, la distribución Chi-cuadrado y la distribución F de Fisher, cada una con aplicaciones específicas y bien definidas en el análisis sanitario.

Estas distribuciones permiten modelar la incertidumbre asociada a estimaciones muestrales, evaluar asociaciones entre variables y comparar dispersiones, constituyendo herramientas esenciales para la inferencia estadística que se desarrollará con mayor profundidad en el capítulo siguiente.

3.6.1 Distribución t de Student

La distribución t de Student se utiliza cuando se desea inferir sobre la media poblacional a partir de muestras pequeñas y cuando la desviación estándar poblacional es desconocida. En estas condiciones, la variabilidad de la estimación es mayor que en el caso de muestras grandes, y la distribución t ajusta este incremento de incertidumbre.

La distribución t depende de un parámetro denominado grados de libertad, generalmente igual a $n - 1$, donde n es el tamaño de la muestra. A medida que

los grados de libertad aumentan, la distribución t se aproxima progresivamente a la distribución normal.

La estadística t se calcula como:

$$t = \frac{\bar{x} - \mu}{s/\sqrt{n}}$$

donde $\bar{x}$ es la media muestral, μ la media poblacional hipotética, s la desviación estándar muestral y n el tamaño de la muestra.

Tabla 3-11. Uso clínico de la distribución t de Student

Situación	Aplicación
Muestra pequeña	Estimación de la media
Varianza desconocida	Comparación con valor de referencia
Estudios clínicos	Evaluación de tratamientos

Nota. Aplicaciones frecuentes de la distribución t en estudios biomédicos con muestras pequeñas.

3.6.2 Distribución Chi-cuadrado

La distribución Chi-cuadrado se emplea principalmente para analizar la asociación entre variables cualitativas y para evaluar la concordancia entre frecuencias observadas y esperadas. En ciencias de la salud, su uso es común en estudios epidemiológicos, análisis de tablas de contingencia y evaluación de pruebas diagnósticas.

La estadística Chi-cuadrado se define como:

$$\chi^2 = \sum \frac{(O-E)^2}{E}$$

donde O representa las frecuencias observadas y E las frecuencias esperadas bajo la hipótesis de independencia.

Esta distribución depende de los grados de libertad, que se calculan en función del número de categorías analizadas. La interpretación sanitaria se centra en

determinar si las diferencias entre frecuencias observadas y esperadas pueden atribuirse al azar o reflejan una asociación real (Pagano y Gauvreau, 2018).

Tabla 3-12. Aplicaciones sanitarias de la distribución Chi-cuadrado

Contexto	Ejemplo
Epidemiología	Asociación entre factor y enfermedad
Diagnóstico	Concordancia de resultados
Salud pública	Comparación de proporciones

Nota. Aplicaciones sanitarias donde la distribución Chi-cuadrado permite analizar asociaciones entre variables categóricas.

3.6.3 Distribución F de Fisher

La distribución F de Fisher se utiliza para comparar varianzas y evaluar la igualdad de dispersión entre dos o más grupos. En investigación biomédica, esta distribución es fundamental en el análisis de varianza y en la comparación de precisión entre métodos de medición.

La estadística F se expresa como el cociente entre dos varianzas:

$$F = \frac{s_1^2}{s_2^2}$$

donde s_1^2 y s_2^2 corresponden a las varianzas de dos grupos independientes. La distribución F depende de dos conjuntos de grados de libertad asociados a cada varianza.

Tabla 3-13. Uso clínico de la distribución F de Fisher

Aplicación	Interpretación
Comparación de varianzas	Diferencias en dispersión
Análisis de tratamientos	Variabilidad entre grupos
Control de calidad	Precisión analítica

Nota. Aplicaciones biomédicas donde la distribución F permite evaluar diferencias de variabilidad.

3.6.4 Relación entre las distribuciones y la práctica sanitaria

Cada una de estas distribuciones responde a una necesidad específica en el análisis sanitario. La distribución t de Student se orienta a la estimación de medias con muestras pequeñas, la Chi-cuadrado a la evaluación de asociaciones categóricas y la F de Fisher a la comparación de variabilidades. Su uso adecuado permite seleccionar métodos estadísticos coherentes con el tipo de datos y el diseño del estudio (Daniel y Cross, 2017).

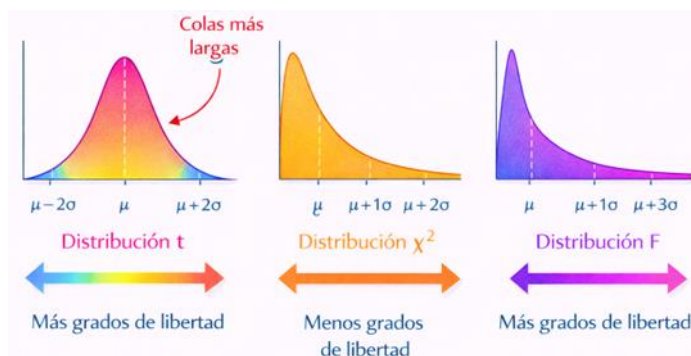


Figura 3-6. Distribuciones t, Chi-cuadrado y F

Ejercicio: identificación de la distribución adecuada

Planteamiento

Se presentan las siguientes situaciones:

1. Comparar la media de colesterol de un grupo pequeño con un valor de referencia.
2. Evaluar la asociación entre sexo y presencia de hipertensión.
3. Comparar la variabilidad de dos métodos de medición de glucosa.

Actividad

Indicar la distribución estadística más adecuada para cada situación.

Solución

1. Distribución t de Student.
2. Distribución Chi-cuadrado.
3. Distribución F de Fisher.

Interpretación

La correcta selección de la distribución estadística garantiza que el análisis sea coherente con la naturaleza de los datos y la pregunta sanitaria planteada, fortaleciendo la validez de las conclusiones.

3.7 Identificación de valores atípicos y control de calidad en mediciones

La identificación de valores atípicos y el control de calidad en las mediciones constituyen componentes esenciales de la bioestadística aplicada a las ciencias de la salud. Los datos biomédicos están sujetos a variabilidad biológica, errores de medición y fallas en los procedimientos analíticos, lo que puede generar observaciones que se apartan de manera notable del patrón general. Reconocer estos valores y gestionarlos adecuadamente resulta indispensable para garantizar la validez de los análisis estadísticos y la seguridad de las decisiones clínicas basadas en ellos.

En la práctica sanitaria, un valor atípico no debe interpretarse automáticamente como un error. Puede representar un evento clínico real de importancia diagnóstica o pronóstica. Por ello, la bioestadística aplicada propone criterios objetivos para su identificación y estrategias sistemáticas de control de calidad que permitan distinguir entre variabilidad esperable y anomalías metodológicas.

3.7.1 Concepto de valor atípico en datos biomédicos

Un valor atípico se define como una observación que difiere de forma marcada del resto de los datos de una muestra. En mediciones clínicas o de laboratorio, estos valores pueden surgir por errores preanalíticos, analíticos o postanalíticos, así como por condiciones fisiológicas o patológicas poco frecuentes en la población estudiada.

Desde el punto de vista estadístico, los valores atípicos alteran medidas como la media y la desviación estándar, afectando la interpretación global de los resultados. Por esta razón, su identificación temprana forma parte del análisis descriptivo inicial en cualquier estudio sanitario.

3.7.2 Métodos estadísticos para la identificación de valores atípicos

Existen diversos métodos estadísticos para detectar valores atípicos, cuya selección depende del tipo de datos y del tamaño muestral. Entre los más utilizados en ciencias de la salud se encuentran los basados en la distribución normal y los basados en medidas de posición.

Uno de los criterios más empleados es el uso de la puntuación *z*. Un valor se considera atípico cuando su puntuación *z* supera un umbral convencional, generalmente ± 3 :

$$z = \frac{x - \mu}{\sigma}$$

Cuando $|z| > 3$, el valor se considera inusual dentro de una población que sigue aproximadamente una distribución normal.

Tabla 3-14. Criterios estadísticos para identificar valores atípicos

Método	Criterio
Puntuación <i>z</i>	(
Rango intercuartílico	Valores fuera de $Q1-1.5 \cdot RIQ$ o $Q3+1.5 \cdot RIQ$
Gráficos exploratorios	Observaciones extremas

Nota. Criterios estadísticos habituales para detectar valores atípicos en datos clínicos y de laboratorio.

3.7.3 Representación gráfica de valores atípicos

Las herramientas gráficas constituyen un complemento esencial para la identificación de valores atípicos. Diagramas de caja permiten visualizar la mediana, los cuartiles y los valores extremos de manera clara, facilitando la

detección de observaciones inusuales sin depender exclusivamente de cálculos numéricos.

En el laboratorio clínico, estas representaciones gráficas se utilizan de forma rutinaria para evaluar la estabilidad de los procesos analíticos y detectar desviaciones que requieran investigación adicional.



Figura 3-7. Identificación gráfica de valores atípicos

3.7.4 Control de calidad en mediciones sanitarias

El control de calidad en mediciones biomédicas tiene como objetivo asegurar que los resultados obtenidos sean precisos, reproducibles y clínicamente confiables. En el laboratorio clínico, este control se aplica de manera sistemática para detectar errores analíticos y garantizar la comparabilidad de los resultados a lo largo del tiempo.

Desde el punto de vista estadístico, el control de calidad se apoya en el seguimiento continuo de indicadores como la media, la desviación estándar y la variabilidad de los resultados obtenidos en muestras de control.

Tabla 3-15. Componentes del control de calidad en mediciones sanitarias

Componente	Finalidad
Control interno	Monitoreo diario del proceso
Control externo	Comparación interlaboratorios
Análisis estadístico	Detección de desviaciones

Nota. Elementos básicos del control de calidad aplicados a mediciones clínicas y de laboratorio.

3.7.5 Relación entre valores atípicos y control de calidad

La identificación de valores atípicos y el control de calidad están estrechamente vinculados. Un aumento repentino de valores extremos puede indicar problemas en el proceso de medición, como fallas en equipos, reactivos o procedimientos. Por el contrario, la ausencia total de variabilidad puede sugerir errores de registro o manipulación de datos.

La bioestadística aplicada recomienda investigar el origen de los valores atípicos antes de decidir su exclusión, considerando tanto criterios estadísticos como clínicos. Esta práctica protege la integridad de los datos y evita sesgos analíticos que puedan comprometer las conclusiones sanitarias.

Ejercicio: detección de valores atípicos en datos de laboratorio

Planteamiento

Se registran los siguientes valores de glucosa en ayunas (mg/dL) en diez pacientes:

82, 85, 88, 90, 92, 95, 98, 100, 105, 160.

Actividad

1. Calcular la media y la desviación estándar.
2. Determinar la puntuación z del valor 160 mg/dL.
3. Identificar si se trata de un valor atípico.
4. Interpretar el resultado desde el punto de vista clínico.

Solución

Media aproximada

$$\mu = \frac{82 + 85 + 88 + 90 + 92 + 95 + 98 + 100 + 105 + 160}{10} = 99.5$$

Desviación estándar aproximada

$$\sigma = 22.4$$

Puntuación z del valor 160

$$z = \frac{160 - 99.5}{22.4} = 2.7$$

El valor de glucosa de 160 mg/dL presenta una puntuación z elevada, cercana al umbral de valor atípico. Desde el punto de vista clínico, este resultado requiere verificación analítica y evaluación del estado metabólico del paciente antes de ser descartado como error.

Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

CAPÍTULO IV

Inferencia Estadística en la Investigación Sanitaria

AUTOR: Hidalgo Villavicencio Gilson Alfonso



4 Inferencia Estadística en la Investigación Sanitaria

4.1 Conceptos de estimación puntual e intervalos de confianza

La inferencia estadística constituye el puente metodológico entre los datos observados en una muestra y las conclusiones que se desean establecer sobre una población de interés. En investigación sanitaria, este proceso resulta esencial, ya que las decisiones clínicas, epidemiológicas y de salud pública rara vez se basan en el estudio exhaustivo de toda la población, sino en información obtenida a partir de muestras representativas. Dentro de este marco, la estimación puntual y los intervalos de confianza representan herramientas fundamentales para cuantificar parámetros poblacionales y expresar la incertidumbre inherente a dichas estimaciones.

A diferencia de la estadística descriptiva, que se limita a resumir datos observados, la inferencia estadística permite formular conclusiones más allá de la muestra, siempre bajo supuestos claramente definidos y con un nivel de incertidumbre cuantificable. Este enfoque es central en la investigación clínica y biomédica, donde los resultados deben interpretarse con rigor y prudencia.

4.1.1 Concepto de estimación puntual

La estimación puntual consiste en utilizar un estadístico muestral para aproximar el valor de un parámetro poblacional desconocido. En ciencias de la salud, los parámetros de interés más frecuentes incluyen la media poblacional, la proporción de individuos con una determinada condición y la varianza de una medición biomédica.

Por ejemplo, la media muestral $\bar{x}$ se utiliza como estimador puntual de la media poblacional μ , mientras que la proporción muestral $\hat{p}$ se emplea como estimador puntual de la proporción poblacional p . Estos estimadores se calculan directamente a partir de los datos observados y proporcionan una única cifra como aproximación del parámetro real.

Desde el punto de vista sanitario, la estimación puntual resulta útil para describir de forma concisa un fenómeno, como el valor promedio de una variable clínica o la frecuencia observada de una enfermedad en una muestra. Sin embargo, esta estimación no informa sobre la precisión del estimador ni sobre la variabilidad asociada al proceso de muestreo (Motulsky, 2015).

4.1.2 Propiedades de los estimadores en salud

En bioestadística aplicada, un buen estimador puntual debe cumplir ciertas propiedades deseables. Entre las más relevantes se encuentran la insesgadez, la consistencia y la eficiencia. Un estimador insesgado presenta un valor esperado igual al parámetro poblacional, mientras que un estimador consistente se aproxima al valor real a medida que aumenta el tamaño muestral (Rosner, 2016).

En investigación sanitaria, estas propiedades cobran especial importancia, ya que estimaciones sesgadas o inestables pueden conducir a conclusiones clínicas erróneas y afectar la toma de decisiones terapéuticas o preventivas.

Tabla 4-1. Estimadores puntuales habituales en investigación sanitaria

Parámetro poblacional	Estimador muestral	Aplicación clínica
Media μ	Media muestral $\bar{x}$	Valores promedio de laboratorio
Proporción p	Proporción muestral $\hat{p}$	Prevalencia de enfermedad
Varianza σ^2	Varianza muestral s^2	Variabilidad biológica

Nota. Estimadores puntuales utilizados para aproximar parámetros poblacionales en estudios clínicos y epidemiológicos.

4.1.3 Limitaciones de la estimación puntual

Aunque la estimación puntual es sencilla y directa, presenta una limitación fundamental: no refleja la incertidumbre asociada al proceso de muestreo. Dos muestras distintas extraídas de la misma población pueden generar estimaciones puntuales diferentes, aun cuando se haya seguido el mismo diseño metodológico.

En la práctica sanitaria, esta variabilidad puede ser crítica. Por ejemplo, una estimación puntual de la prevalencia de una enfermedad no informa cuán precisa es dicha estimación ni qué tan lejos podría estar del valor real. Para

abordar esta limitación, la inferencia estadística introduce el concepto de intervalo de confianza.

4.1.4 Concepto de intervalo de confianza

Un intervalo de confianza es un rango de valores calculado a partir de los datos muestrales que tiene una probabilidad conocida de contener el verdadero parámetro poblacional. En investigación sanitaria, el nivel de confianza más utilizado es el 95 por ciento, lo que indica que, si se repitiera el procedimiento de muestreo un gran número de veces, aproximadamente el 95 por ciento de los intervalos contruidos contendrían el valor real del parámetro (Rosner, 2016).

Este enfoque permite expresar la estimación acompañada de una medida de precisión, aspecto importante para la interpretación clínica y la toma de decisiones basada en evidencia cuantitativa.

4.1.5 Intervalo de confianza para la media poblacional

Cuando la variable de interés sigue aproximadamente una distribución normal y la desviación estándar poblacional es desconocida, el intervalo de confianza para la media se construye utilizando la distribución t de Student. La expresión general es:

$$IC_{95\%} = \bar{x} \pm t_{\alpha/2, n-1} \frac{s}{\sqrt{n}}$$

donde:

- $\bar{x}$ es la media muestral,
- s es la desviación estándar muestral,
- n es el tamaño de la muestra,
- $t_{\alpha/2, n-1}$ es el valor crítico de la distribución t con $n - 1$ grados de libertad.

Este intervalo se interpreta como el rango plausible de valores para la media poblacional, considerando la variabilidad observada en la muestra (Pagano y Gauvreau, 2018).

Tabla 4-2. Factores que influyen en la amplitud del intervalo de confianza

Factor	Efecto
Tamaño muestral	Mayor tamaño reduce amplitud
Variabilidad	Mayor dispersión amplía el intervalo
Nivel de confianza	Mayor confianza amplía el intervalo

Nota. Elementos que determinan la precisión de los intervalos de confianza en estudios sanitarios.

4.1.6 Interpretación sanitaria de los intervalos de confianza

La interpretación correcta de un intervalo de confianza es esencial en la investigación clínica. Un intervalo estrecho indica una estimación precisa, mientras que un intervalo amplio refleja mayor incertidumbre. En la práctica sanitaria, esta información resulta más informativa que una estimación puntual aislada, ya que permite evaluar la robustez de los resultados y su relevancia clínica (Altman, 1991; Motulsky, 2015).

Es importante destacar que el intervalo de confianza no implica que el parámetro poblacional tenga una probabilidad específica de encontrarse dentro del intervalo una vez calculado, sino que describe la fiabilidad del procedimiento de estimación utilizado.



Figura 4-1. Estimación puntual e intervalo de confianza

Ejercicio: estimación puntual e intervalo de confianza en salud

Planteamiento

En un estudio clínico se mide la presión arterial sistólica en una muestra de 25 pacientes, obteniéndose una media de 130 mmHg y una desviación estándar de 10 mmHg.

Actividad

1. Identificar la estimación puntual de la presión arterial sistólica media.
2. Calcular el intervalo de confianza del 95 por ciento para la media poblacional.
3. Interpretar los resultados desde el punto de vista clínico.

Solución

Estimación puntual

$$\bar{x} = 130 \text{ mmHg}$$

Intervalo de confianza

Para $n = 25$, el valor crítico aproximado es $t_{0.025,24} = 2.064$.

$$IC_{95\%} = 130 \pm 2.064 \left(\frac{10}{\sqrt{25}} \right)$$
$$IC_{95\%} = 130 \pm 4.13 = [125.87, 134.13]$$

Interpretación clínica

El intervalo de confianza indica que la media poblacional de la presión arterial sistólica se encuentra, con alta fiabilidad, entre 125.9 y 134.1 mmHg. Desde el punto de vista clínico, este rango sugiere valores compatibles con hipertensión leve, lo que justifica seguimiento y evaluación adicional.

4.2 Pruebas de hipótesis: formulación y tipos de error

La prueba de hipótesis es un procedimiento central de la inferencia estadística que permite evaluar si los datos observados en una muestra contienen suficiente evidencia para respaldar una afirmación sobre un parámetro

poblacional. En investigación sanitaria, las pruebas de hipótesis se emplean para contrastar efectos terapéuticos, diferencias entre grupos biológicos, asociaciones entre factores de riesgo y resultados clínicos, entre otros objetivos. El proceso se basa en la comparación entre una hipótesis nula, que representa un estado de ausencia de efecto o diferencia, y una hipótesis alternativa, que indica la presencia del efecto o diferencia que se desea detectar (p. ej., tratamiento eficaz, mayor prevalencia, efecto significativo).

Este método asume que la evidencia de la muestra se cuantifica mediante un estadístico de prueba, y la decisión de aceptar o rechazar una hipótesis se basa en criterios previamente definidos como el nivel de significación y la distribución del estadístico bajo la hipótesis nula. Debido a la naturaleza probabilística del muestreo, las decisiones pueden estar sujetas a errores, cuya comprensión es indispensable para interpretar resultados clínicos y tomar decisiones fundamentadas.

4.2.1 Formulación de hipótesis estadísticas

La prueba de hipótesis comienza con la formulación explícita de dos afirmaciones complementarias:

- **Hipótesis nula (H_0):** afirmación que se somete a prueba estadística y que generalmente representa ausencia de efecto, igualdad de parámetros o status quo en la práctica sanitaria.
- **Hipótesis alternativa (H_1 o H_a):** afirmación contraria a la hipótesis nula, que representa la existencia de efecto, diferencia o asociación de interés clínico.

Estas dos hipótesis son mutuamente exclusivas y abarcan todas las posibilidades plausibles para el parámetro de interés. La prueba estadística evalúa si los datos recolectados proporcionan evidencia suficiente para rechazar H_0 en favor de H_1 .

4.2.2 Criterios de decisión: nivel de significación y valor p

Antes de analizar los datos, se establece un nivel de significación α , que representa la tolerancia a cometer un tipo específico de error (error tipo I, que se describirá más adelante). Comúnmente se utiliza $\alpha = 0,05$ o $\alpha = 0,01$, lo

cual implica aceptar un 5 % o 1 % de probabilidad de rechazar injustamente la hipótesis nula. Si el estadístico de prueba calculado a partir de la muestra produce un valor que cae en la región crítica asociada a α , se rechaza H_0 ; de lo contrario, no se rechaza.

El valor p es otra herramienta complementaria: se define como la probabilidad de observar un estadístico de prueba al menos tan extremo como el observado, asumiendo que H_0 es verdadera. Si el valor p es menor o igual que α , se rechaza H_0 ; si es mayor, no se rechaza.

4.2.3 Tipos de error en la prueba de hipótesis

La toma de decisiones en pruebas de hipótesis está sujeta a incertidumbre debido al muestreo aleatorio. Esto puede dar lugar a dos tipos fundamentales de errores de decisión estadística (Type I y Type II), cuya identificación y manejo son críticos en la investigación sanitaria para evitar conclusiones equivocadas que puedan comprometer intervenciones o diagnósticos clínicos.

a) Error de tipo I

Se produce cuando la hipótesis nula es verdadera, pero los datos muestrales llevan a rechazarla. Este error se asocia con el concepto de falso positivo, es decir, concluir que existe un efecto cuando en realidad no lo hay. La probabilidad de cometer un error de tipo I se denota como α , y su valor está predefinido como el nivel de significación de la prueba. Por ejemplo, si $\alpha = 0,05$, existe una probabilidad del 5 % de rechazar H_0 siendo ésta verdadera.

Este tipo de error es especialmente relevante en ensayos clínicos, donde una conclusión incorrecta de efectividad puede conducir a la adopción de un tratamiento que en realidad no aporta beneficio e incluso puede tener efectos adversos, generando costos y daños innecesarios.

b) Error de tipo II

Ocurre cuando la hipótesis nula es falsa, pero no se rechaza, lo cual se interpreta como un falso negativo. En este caso, se falla al detectar un efecto o diferencia que realmente existe. La probabilidad de cometer un error de tipo II se denota como β , y está directamente relacionada con la potencia estadística

de la prueba ($1 - \beta$). Una potencia elevada implica menor riesgo de error tipo II y mayor capacidad de la prueba para detectar efectos reales.

Este error es crítico en estudios clínicos que buscan demostrar eficacia terapéutica. Por ejemplo, si un tratamiento realmente eficaz no es detectado por una prueba con baja potencia, se puede desaprovechar una intervención beneficiosa para la salud pública.

Tabla 4-3. Esquema de decisiones y tipos de error

Decisión estadística	Verdad poblacional	Resultado	Tipo de error
Rechazar H_0	H_0 verdadera	Se concluye efecto inexistente	Error tipo I (α)
No rechazar H_0	H_0 falsa	Se omite detectar efecto real	Error tipo II (β)
Rechazar H_0	H_0 falsa	Correcto	—
No rechazar H_0	H_0 verdadera	Correcto	—

Nota. Relación entre decisiones estadísticas, realidad poblacional y tipos de error en contraste de hipótesis.

4.2.4 Relación entre niveles de significación y errores

Existe un equilibrio metodológico entre los niveles de significación α y la probabilidad de error de tipo II β . Disminuir α para hacer una prueba más estricta frecuentemente incrementa β , elevando la probabilidad de pasar por alto efectos reales. Este intercambio debe considerarse cuidadosamente en el diseño de estudios clínicos, especialmente cuando los costos clínicos de cada tipo de error son distintos.

La potencia estadística también depende del tamaño muestral y del efecto esperado. Estudios con muestras pequeñas suelen tener menor potencia y mayor probabilidad de error tipo II, incluso si el efecto clínico es importante. Por tanto, la planificación del tamaño de la muestra antes de realizar el estudio es una etapa crítica en la investigación sanitaria.

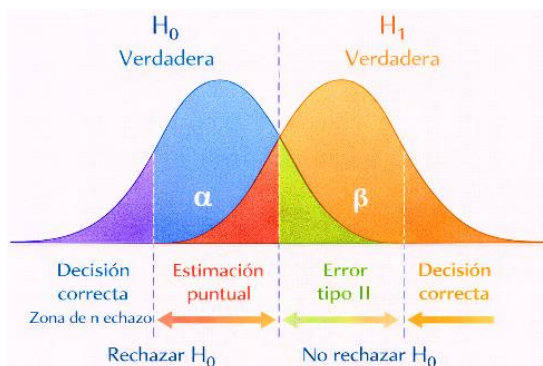


Figura 4-2. Interpretación de errores tipo I y II

Ejercicio: formulación de hipótesis y tipos de error

Planteamiento

En un estudio clínico se quiere evaluar si un nuevo medicamento reduce la presión arterial media más que el tratamiento estándar. Se establece:

- $H_0: \mu_{\text{nuevo}} = \mu_{\text{estándar}}$
- $H_1: \mu_{\text{nuevo}} < \mu_{\text{estándar}}$

Se fija un nivel de significación $\alpha = 0,05$.

Actividad

1. Identificar qué representa un error de tipo I en este contexto.
2. Identificar qué representa un error de tipo II.
3. Interpretar las implicaciones clínicas de cada error.

Solución

1. **Error de tipo I:** rechazar H_0 cuando en realidad las medias son iguales, lo que implicaría concluir que el medicamento nuevo reduce la presión arterial cuando en realidad no lo hace. Esto podría conducir a su adopción clínicamente inapropiada.

2. **Error de tipo II:** no rechazar H_0 cuando en realidad el nuevo medicamento es más eficaz que el estándar, lo que implicaría dejar de recomendar una intervención eficaz, afectando negativamente resultados clínicos y oportunidades terapéuticas.

Interpretación clínica

El error tipo I puede causar que se promueva un tratamiento sin eficacia real, con riesgos y costos adicionales. El error tipo II puede implicar desaprovechar un avance terapéutico real. Balancear estos riesgos mediante diseño adecuado, selección de α y cálculo de potencia es esencial en investigación sanitaria.

4.3 Prueba t para muestras independientes y relacionadas

La prueba t de Student es una de las herramientas más utilizadas en la investigación sanitaria para comparar medias y evaluar diferencias entre grupos o condiciones. Su relevancia radica en que muchas preguntas clínicas se formulan en términos de diferencias promedio, como la efectividad de un tratamiento frente a otro o el cambio observado en un mismo grupo de pacientes antes y después de una intervención. La prueba t permite determinar si las diferencias observadas en las muestras pueden atribuirse al azar o reflejan un efecto real en la población de interés (Altman, 1991; Rosner, 2016).

En el ámbito de las ciencias de la salud, la prueba t se aplica bajo supuestos bien definidos, y su correcta selección depende del diseño del estudio y de la relación existente entre las muestras comparadas. Distinguir entre muestras independientes y muestras relacionadas es fundamental para evitar errores metodológicos y conclusiones inválidas (Daniel y Cross, 2017).

4.3.1 Supuestos generales de la prueba t

Antes de aplicar la prueba t, es necesario verificar una serie de supuestos que garantizan la validez del análisis:

- La variable de interés es cuantitativa y medida en escala de intervalo o razón.
- Las observaciones provienen de poblaciones con distribución aproximadamente normal.

- Las observaciones son independientes dentro de cada grupo, salvo en el caso de muestras relacionadas.
- La varianza poblacional es desconocida y se estima a partir de la muestra.

En estudios clínicos, la normalidad se evalúa mediante métodos gráficos y pruebas estadísticas, especialmente cuando el tamaño muestral es pequeño (Motulsky, 2015).

4.3.2 Prueba t para muestras independientes

La prueba t para muestras independientes se utiliza cuando se comparan las medias de dos grupos distintos que no tienen relación entre sí. Ejemplos típicos incluyen la comparación de un grupo de pacientes que recibe un nuevo tratamiento con otro que recibe el tratamiento estándar, o la comparación de parámetros clínicos entre hombres y mujeres (Pagano y Gauvreau, 2018).

La hipótesis nula plantea que las medias poblacionales de ambos grupos son iguales:

$$H_0: \mu_1 = \mu_2$$

La estadística t se calcula como:

$$t = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

donde $\bar{x}_1$ y $\bar{x}_2$ son las medias muestrales, s_1^2 y s_2^2 las varianzas muestrales, y n_1 y n_2 los tamaños muestrales de cada grupo.

Tabla 4-4. Aplicación clínica de la prueba t para muestras independientes

Situación clínica	Comparación
Ensayo clínico	Tratamiento nuevo vs estándar
Estudios poblacionales	Dos grupos independientes

Nota. Contextos sanitarios donde se aplica la prueba t para comparar medias entre grupos independientes.

4.3.3 Prueba t para muestras relacionadas

La prueba t para muestras relacionadas se emplea cuando las observaciones están emparejadas o corresponden al mismo individuo medido en dos momentos distintos. Este diseño es común en estudios pre y post intervención, donde se evalúa el efecto de un tratamiento en un mismo grupo de pacientes (Altman, 1991).

En este caso, el análisis se centra en las diferencias individuales entre pares de observaciones. La hipótesis nula se formula como:

$$H_0: \mu_d = 0$$

donde μ_d representa la media poblacional de las diferencias.

La estadística t se calcula como:

$$t = \frac{\bar{d}}{s_d / \sqrt{n}}$$

donde $\bar{d}$ es la media de las diferencias, s_d la desviación estándar de las diferencias y n el número de pares.

Tabla 4-5. Uso de la prueba t para muestras relacionadas

Diseño	Ejemplo
Pre y post intervención	Presión arterial antes y después
Estudios longitudinales	Seguimiento clínico
Ensayos cruzados	Misma persona en dos condiciones

Nota. Aplicaciones típicas de la prueba t cuando las observaciones están emparejadas.

4.3.4 Comparación entre pruebas t independientes y relacionadas

La diferencia entre ambas pruebas radica en la relación entre las observaciones. Utilizar una prueba t independiente cuando los datos están emparejados reduce la potencia estadística, mientras que aplicar una prueba relacionada a datos independientes viola los supuestos del método.



Figura 4-3. Selección de la prueba t según el diseño del estudio

Ejercicio: prueba t en investigación sanitaria

Planteamiento

Se evalúa el efecto de un programa de ejercicio en la presión arterial sistólica de 12 pacientes. Los valores promedio antes del programa son 145 mmHg y después 135 mmHg. La desviación estándar de las diferencias es 8 mmHg.

Actividad

1. Identificar el tipo de prueba t adecuada.
2. Calcular la estadística t.
3. Interpretar el resultado clínicamente.

Solución

Tipo de prueba

Muestras relacionadas, ya que se evalúan los mismos pacientes antes y después.

Cálculo

$$\begin{aligned}\bar{d} &= 145 - 135 = 10 \\ t &= \frac{10}{8/\sqrt{12}} = \frac{10}{2.31} = 4.33\end{aligned}$$

Con 11 grados de libertad, este valor es superior al valor crítico habitual para $\alpha = 0.05$.

Interpretación clínica

La diferencia observada en la presión arterial es estadísticamente significativa. Desde el punto de vista clínico, el programa de ejercicio se asocia con una reducción relevante de la presión arterial sistólica en los pacientes evaluados.

4.4 Prueba Chi-cuadrado en análisis de proporciones clínicas

La prueba Chi-cuadrado constituye uno de los procedimientos estadísticos más utilizados en la investigación sanitaria para analizar proporciones y evaluar la asociación entre variables cualitativas. En ciencias de la salud, muchas preguntas de investigación no se centran en valores promedio, sino en la distribución de categorías, como la presencia o ausencia de enfermedad, el resultado positivo o negativo de una prueba, o la pertenencia a distintos grupos clínicos. La prueba Chi-cuadrado permite determinar si las diferencias observadas entre proporciones pueden atribuirse al azar o reflejan una relación real entre las variables estudiadas.

Este procedimiento resulta especialmente útil en estudios epidemiológicos, análisis de resultados diagnósticos y evaluaciones de programas de salud, donde los datos suelen organizarse en tablas de contingencia.

4.4.1 Fundamento de la prueba Chi-cuadrado

La prueba Chi-cuadrado compara las frecuencias observadas en cada categoría con las frecuencias que se esperarían si no existiera asociación entre las variables. La hipótesis nula plantea que las variables son independientes, mientras que la hipótesis alternativa indica que existe algún grado de asociación.

Formalmente, la hipótesis nula se expresa como:

$$H_0: \text{Las variables son independientes}$$

La hipótesis alternativa se formula como:

$$H_1: \text{Las variables no son independientes}$$

4.4.2 Cálculo de la estadística Chi-cuadrado

La estadística Chi-cuadrado se define como:

$$\chi^2 = \sum \frac{(O-E)^2}{E}$$

donde:

- O representa la frecuencia observada en cada celda,
- E representa la frecuencia esperada bajo la hipótesis de independencia.

Las frecuencias esperadas se calculan a partir de los totales marginales de la tabla de contingencia:

$$E = \frac{(\text{Total fila}) \times (\text{Total columna})}{\text{Total general}}$$

Este procedimiento permite cuantificar el grado de discrepancia entre lo observado y lo esperado si no existiera asociación.

Tabla 4-6. Ejemplo de tabla de contingencia en análisis clínico

Diagnóstico	Prueba positiva	Prueba negativa	Total
Enfermedad presente	O ₁	O ₂	Total 1
Enfermedad ausente	O ₃	O ₄	Total 2
Total	Total +	Total –	Total general

Nota. Estructura típica de una tabla de contingencia utilizada para aplicar la prueba Chi-cuadrado en estudios clínicos.

4.4.3 Grados de libertad e interpretación

Los grados de libertad de la prueba Chi-cuadrado dependen del tamaño de la tabla y se calculan como:

$$gl = (r - 1)(c - 1)$$

donde r es el número de filas y c el número de columnas. El valor calculado de χ^2 se compara con un valor crítico obtenido de la distribución Chi-cuadrado para el nivel de significación seleccionado y los grados de libertad correspondientes (Rosner, 2016).

4.4.4 Supuestos y consideraciones en salud

Para aplicar la prueba Chi-cuadrado de manera válida, se deben cumplir ciertos supuestos:

- Las observaciones son independientes.
- Las categorías son mutuamente excluyentes.
- Las frecuencias esperadas no son demasiado pequeñas.

En estudios sanitarios, se recomienda que la mayoría de las frecuencias esperadas sean mayores o iguales a 5 para garantizar una aproximación adecuada de la distribución Chi-cuadrado.

Tabla 4-7. Supuestos de la prueba Chi-cuadrado

Supuesto	Implicación
Independencia	Cada individuo se cuenta una vez
Categorías exclusivas	No superposición
Frecuencias suficientes	Validez del contraste

Nota. Condiciones necesarias para la aplicación adecuada de la prueba Chi-cuadrado en investigación sanitaria.

4.4.5 Interpretación clínica de la prueba Chi-cuadrado

Cuando el resultado de la prueba Chi-cuadrado es estadísticamente significativo, se concluye que existe asociación entre las variables analizadas. Sin embargo, esta prueba no indica la dirección ni la magnitud de la asociación, por lo que debe complementarse con medidas adicionales, como razones de riesgo u odds ratio, especialmente en estudios epidemiológicos.

Desde el punto de vista clínico, identificar asociaciones significativas permite orientar estrategias diagnósticas, preventivas o terapéuticas basadas en patrones observados en la población estudiada.

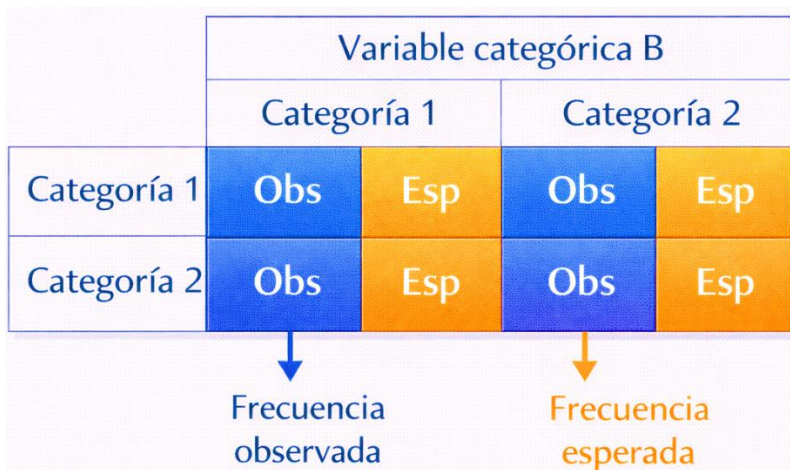


Figura 4-4. Asociación entre variables cualitativas en salud

Ejercicio aplicado: prueba Chi-cuadrado en un estudio clínico

Planteamiento

En un estudio se evalúa la asociación entre tabaquismo y presencia de enfermedad respiratoria en 100 personas, obteniéndose los siguientes datos:

	Enfermedad	No enfermedad	Total
Fumadores	30	20	50
No fumadores	10	40	50
Total	40	60	100

Actividad

1. Calcular las frecuencias esperadas.
2. Calcular la estadística Chi-cuadrado.
3. Interpretar el resultado desde el punto de vista clínico.

Solución

Frecuencias esperadas

$$E_{11} = \frac{50 \times 40}{100} = 20, E_{12} = 30$$
$$E_{21} = 20, E_{22} = 30$$

Cálculo de χ^2

$$\chi^2 = \frac{(30-20)^2}{20} + \frac{(20-30)^2}{30} + \frac{(10-20)^2}{20} + \frac{(40-30)^2}{30}$$
$$\chi^2 = 5 + 3.33 + 5 + 3.33 = 16.66$$

Con 1 grado de libertad, este valor es superior al valor crítico habitual para $\alpha = 0.05$.

Interpretación clínica

Existe una asociación estadísticamente significativa entre tabaquismo y enfermedad respiratoria. Desde el punto de vista clínico y sanitario, este resultado respalda la relación entre el consumo de tabaco y el desarrollo de patología respiratoria, justificando intervenciones preventivas y programas de cesación tabáquica.

4.5 ANOVA en comparación de tratamientos o grupos biológicos

El análisis de varianza, conocido como ANOVA, es un procedimiento estadístico diseñado para comparar las medias de tres o más grupos de manera simultánea. En investigación sanitaria, esta técnica resulta indispensable cuando se desea evaluar la eficacia de diferentes tratamientos, comparar respuestas biológicas entre varios grupos poblacionales o analizar variaciones de parámetros clínicos bajo distintas condiciones experimentales. El ANOVA permite determinar si las diferencias observadas entre las medias de los grupos pueden atribuirse al azar o si reflejan efectos reales asociados a los factores estudiados.

A diferencia de realizar múltiples comparaciones dos a dos mediante pruebas *t*, lo cual incrementa el riesgo de error tipo I, el ANOVA controla este riesgo al evaluar todas las medias en un único contraste global, manteniendo el rigor estadístico del análisis.

4.5.1 Fundamento del análisis de varianza

El principio central del ANOVA consiste en descomponer la variabilidad total observada en los datos en dos componentes: la variabilidad entre los grupos y la variabilidad dentro de los grupos. Si las diferencias entre las medias de los grupos son grandes en relación con la variabilidad interna, se concluye que al menos una de las medias difiere significativamente de las demás.

La hipótesis nula del ANOVA se formula como:

$$H_0: \mu_1 = \mu_2 = \mu_3 = \dots = \mu_k$$

donde μ_i representa la media poblacional del grupo *i*. La hipótesis alternativa establece que al menos una media difiere del resto.

4.5.2 Estadística F y su interpretación

La estadística de contraste del ANOVA es la estadística F, que se calcula como el cociente entre la variabilidad entre grupos y la variabilidad dentro de los grupos:

$$F = \frac{\text{Variabilidad entre grupos}}{\text{Variabilidad dentro de los grupos}} = \frac{MS_{entre}}{MS_{dentro}}$$

donde MS representa los cuadrados medios obtenidos a partir de las sumas de cuadrados correspondientes. Un valor elevado de F indica que las diferencias entre las medias de los grupos son mayores de lo que cabría esperar por azar.

Tabla 4-8. Componentes del análisis de varianza

Componente	Descripción
Variabilidad entre grupos	Diferencias entre medias
Variabilidad dentro de grupos	Dispersión interna
Estadística F	Relación entre variabilidades

Nota. Elementos básicos que intervienen en el cálculo y la interpretación del análisis de varianza.

4.5.3 Supuestos del ANOVA en estudios sanitarios

Para que los resultados del ANOVA sean válidos, deben cumplirse ciertos supuestos:

- Las observaciones son independientes.
- Las variables analizadas son cuantitativas.
- Las poblaciones presentan distribución aproximadamente normal.
- Las varianzas de los grupos son similares.

En investigación biomédica, estos supuestos se evalúan mediante análisis descriptivos previos y pruebas específicas cuando es necesario.

4.5.4 Interpretación clínica del ANOVA

Cuando el resultado del ANOVA es estadísticamente significativo, se concluye que existen diferencias entre las medias de los grupos analizados. Sin embargo, el ANOVA no indica cuáles grupos difieren entre sí. Para identificar estas diferencias específicas, se requieren análisis posteriores de comparación múltiple, que permiten localizar los contrastes responsables del efecto global.

Desde el punto de vista sanitario, este enfoque es especialmente útil en ensayos clínicos con múltiples tratamientos, donde se desea evaluar cuál intervención produce una respuesta clínica más favorable.

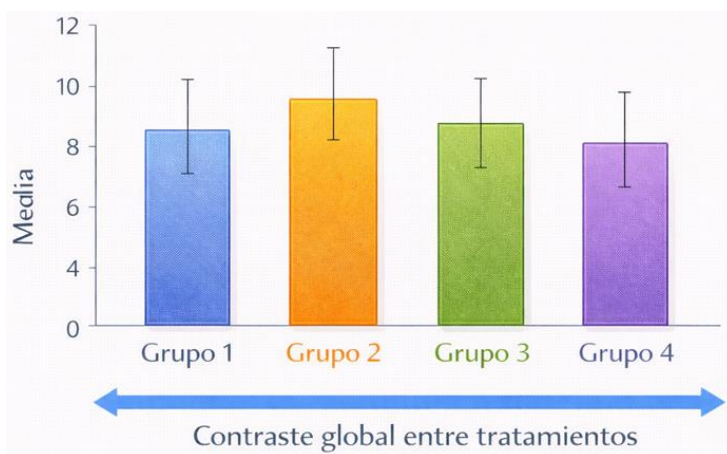


Figura 4-5. Comparación de medias mediante ANOVA

Ejercicio aplicado: ANOVA en un estudio clínico

Planteamiento

Se evalúa el efecto de tres tratamientos antihipertensivos sobre la presión arterial sistólica. Se obtienen las siguientes medias en mmHg:

- Tratamiento A: 138
- Tratamiento B: 130
- Tratamiento C: 125

Cada grupo está compuesto por 10 pacientes, y la variabilidad dentro de los grupos es similar.

Actividad

1. Formular la hipótesis nula del ANOVA.
2. Explicar el criterio de decisión basado en la estadística F.
3. Interpretar el resultado desde el punto de vista clínico.

Solución

Hipótesis nula

$$H_0: \mu_A = \mu_B = \mu_C$$

Criterio de decisión

Si el valor calculado de la estadística F supera el valor crítico correspondiente al nivel de significación seleccionado, se rechaza la hipótesis nula.

Interpretación clínica

Un resultado significativo indicaría que al menos uno de los tratamientos produce una reducción de la presión arterial diferente a los demás. Este hallazgo justificaría análisis posteriores para identificar el tratamiento más efectivo.

4.6 Correlación y regresión lineal simple

En la investigación sanitaria es frecuente analizar la relación entre dos variables cuantitativas, como la asociación entre edad y presión arterial, índice de masa corporal y glucosa sanguínea, o dosis de un fármaco y respuesta clínica. La correlación y la regresión lineal simple constituyen herramientas estadísticas fundamentales para describir y cuantificar este tipo de relaciones, permitiendo evaluar la intensidad, dirección y forma de la asociación entre variables, así como realizar estimaciones predictivas bajo supuestos claramente definidos.

Mientras la correlación se centra en medir el grado de asociación entre dos variables, la regresión lineal simple permite modelar una relación funcional,

en la que una variable se considera predictora de otra. Ambas técnicas son ampliamente utilizadas en estudios clínicos, epidemiológicos y de laboratorio.

4.6.1 Concepto de correlación en salud

La correlación cuantifica la fuerza y la dirección de la relación lineal entre dos variables cuantitativas. El coeficiente de correlación más utilizado en investigación sanitaria es el coeficiente de correlación de Pearson, que toma valores entre -1 y $+1$. Un valor cercano a $+1$ indica una relación positiva fuerte, un valor cercano a -1 indica una relación negativa fuerte, y un valor cercano a 0 indica ausencia de relación lineal.

La expresión matemática del coeficiente de correlación de Pearson es:

$$r = \frac{\sum(x - \bar{x})(y - \bar{y})}{\sqrt{\sum(x - \bar{x})^2 \sum(y - \bar{y})^2}}$$

Este coeficiente es adimensional y no depende de las unidades de medida, lo que facilita la comparación entre distintos estudios biomédicos.

Tabla 4-9. Interpretación general del coeficiente de correlación

Valor de r	Interpretación
0.00 a 0.19	Relación muy débil
0.20 a 0.39	Relación débil
0.40 a 0.59	Relación moderada
0.60 a 0.79	Relación fuerte
0.80 a 1.00	Relación muy fuerte

Nota. Guía orientativa para interpretar la magnitud de la correlación lineal entre variables cuantitativas en salud.

4.6.2 Limitaciones de la correlación

Es fundamental destacar que la correlación no implica causalidad. Una asociación estadística entre dos variables no significa que una sea la causa de la otra. En investigación sanitaria, correlaciones espurias pueden surgir por la

influencia de variables de confusión o por coincidencias asociadas al tamaño muestral.

Por esta razón, la correlación debe interpretarse siempre en conjunto con el diseño del estudio, el conocimiento biológico y la plausibilidad clínica.

4.6.3 Regresión lineal simple

La regresión lineal simple permite describir la relación entre una variable dependiente Y y una variable independiente X mediante una ecuación lineal. En el contexto sanitario, esta técnica se utiliza para modelar cómo cambia una variable clínica en función de otra, como el aumento de la presión arterial con la edad o la variación de la glucemia con el índice de masa corporal.

La ecuación general del modelo de regresión lineal simple es:

$$Y = \beta_0 + \beta_1 X + \varepsilon$$

donde:

- β_0 es la ordenada al origen,
- β_1 es la pendiente, que indica el cambio promedio de Y por unidad de cambio en X ,
- ε representa el error aleatorio.

4.6.4 Interpretación clínica de la pendiente

La pendiente del modelo de regresión tiene una interpretación directa en términos clínicos. Por ejemplo, una pendiente positiva indica que el valor esperado de la variable dependiente aumenta a medida que aumenta la variable independiente. Esta interpretación facilita la traducción de resultados estadísticos a conclusiones comprensibles para la práctica clínica (Altman, 1991).

Tabla 4-10. Componentes del modelo de regresión lineal simple

Componente	Interpretación
------------	----------------

β_0	Valor esperado de Y cuando X es cero
β_1	Cambio promedio de Y por unidad de X
Error	Variabilidad no explicada

Nota. Elementos del modelo de regresión lineal simple y su significado en investigación sanitaria.

4.6.5 Representación gráfica de la regresión

La relación entre dos variables se visualiza mediante un diagrama de dispersión, sobre el cual se ajusta la recta de regresión. Este gráfico permite evaluar la adecuación del modelo lineal, identificar patrones no lineales y detectar valores atípicos que puedan influir en el análisis.

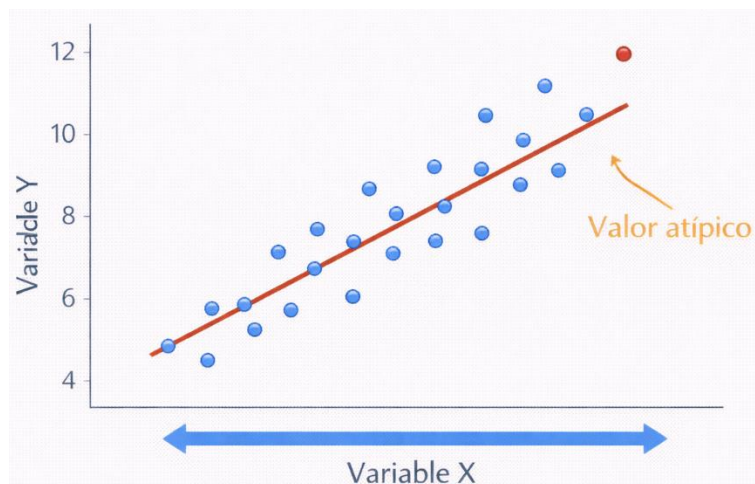


Figura 4-6. Correlación y regresión lineal en datos clínicos

Ejercicio aplicado: correlación y regresión en salud

Planteamiento

En un estudio se analizan los datos de 10 pacientes, registrándose la edad (años) y la presión arterial sistólica (mmHg):

Edad: 40, 45, 50, 55, 60, 65, 70, 75, 80, 85

Presión: 120, 122, 125, 128, 132, 135, 138, 142, 145, 148

Actividad

1. Calcular e interpretar la correlación entre edad y presión arterial.
2. Ajustar un modelo de regresión lineal simple.
3. Interpretar los resultados desde el punto de vista clínico.

Solución conceptual

El coeficiente de correlación es positivo y elevado, lo que indica una relación directa entre edad y presión arterial. El modelo de regresión muestra una pendiente positiva, sugiriendo que la presión arterial sistólica aumenta de forma progresiva con la edad.

Interpretación clínica

Estos resultados son coherentes con el conocimiento fisiológico sobre el envejecimiento vascular. Desde el punto de vista sanitario, la regresión lineal permite estimar el incremento promedio de la presión arterial asociado al aumento de la edad, apoyando estrategias de seguimiento y prevención en poblaciones de mayor riesgo.

4.7 Interpretación de resultados estadísticos en estudios clínicos

La interpretación adecuada de los resultados estadísticos constituye una etapa decisiva en la investigación sanitaria, ya que de ella dependen las conclusiones clínicas, las recomendaciones terapéuticas y las decisiones en salud pública. Un análisis estadístico correcto pierde valor si los resultados no se interpretan con rigor metodológico y sentido clínico. En este contexto, la bioestadística aplicada no se limita al cálculo de valores p o intervalos de confianza, sino que exige integrar la evidencia cuantitativa con el conocimiento biomédico, el diseño del estudio y las implicaciones prácticas de los hallazgos.

En estudios clínicos, la interpretación estadística debe responder a preguntas como la relevancia clínica de los resultados, la precisión de las estimaciones y la validez de las conclusiones inferidas a partir de la muestra analizada.

4.7.1 Significación estadística y relevancia clínica

Uno de los errores más frecuentes en investigación sanitaria consiste en equiparar significación estadística con relevancia clínica. Un resultado puede ser estadísticamente significativo y, sin embargo, carecer de importancia clínica si el tamaño del efecto es pequeño o irrelevante para la práctica médica.

Por ejemplo, una reducción mínima en un parámetro clínico puede alcanzar significación estadística en muestras grandes, pero no justificar un cambio en la conducta terapéutica. Por ello, se recomienda evaluar siempre el tamaño del efecto y su impacto clínico, además del valor p .

4.7.2 Intervalos de confianza como herramienta interpretativa

Los intervalos de confianza proporcionan información más rica que el valor p aislado, ya que permiten evaluar tanto la dirección como la magnitud del efecto, así como la precisión de la estimación. En estudios clínicos, un intervalo estrecho indica alta precisión, mientras que un intervalo amplio refleja mayor incertidumbre.

Además, la posición del intervalo respecto a valores clínicamente relevantes permite juzgar si el efecto observado tiene importancia práctica. Por ejemplo, un intervalo de confianza que excluye un valor clínicamente neutro sugiere un efecto consistente y potencialmente relevante.

Tabla 4-11. Elementos para interpretar resultados estadísticos

Elemento	Pregunta orientadora
Valor p	¿Existe evidencia estadística?
Intervalo de confianza	¿Qué tan preciso es el resultado?
Tamaño del efecto	¿Es clínicamente relevante?
Diseño del estudio	¿Es válida la inferencia?

Nota. Aspectos fundamentales que deben considerarse al interpretar resultados estadísticos en investigación clínica.

4.7.3 Consideración de errores y potencia estadística

La interpretación de resultados estadísticos debe contemplar la posibilidad de errores de tipo I y tipo II. Un resultado no significativo puede deberse a la ausencia real de efecto o a una potencia estadística insuficiente. En estudios con tamaño muestral reducido, la falta de significación no debe interpretarse automáticamente como ausencia de efecto clínico.

En investigación sanitaria, se recomienda reportar y discutir la potencia del estudio, especialmente cuando los resultados no son estadísticamente significativos, para evitar conclusiones erróneas o excesivamente categóricas.

4.7.4 Integración del análisis estadístico con el contexto clínico

La interpretación final de los resultados debe integrar la evidencia estadística con el contexto clínico y biológico. Factores como la plausibilidad fisiopatológica, la consistencia con estudios previos y la calidad metodológica del diseño influyen en la solidez de las conclusiones.

Un análisis estadístico aislado, sin referencia al contexto clínico, puede conducir a interpretaciones simplistas o equivocadas. La bioestadística aplicada promueve una visión integradora, donde los números respaldan, pero no sustituyen, el razonamiento clínico.



Figura 4-7. Proceso de interpretación de resultados estadísticos en estudios clínicos

Ejercicio aplicado: interpretación integral de resultados

Planteamiento

Un ensayo clínico compara un nuevo fármaco con un tratamiento estándar en la reducción de colesterol LDL. El resultado muestra una diferencia media de -8 mg/dL, con un intervalo de confianza del 95 por ciento de -15 a -1 mg/dL y un valor p de 0.03 .

Actividad

1. Determinar si el resultado es estadísticamente significativo.
2. Analizar la precisión de la estimación.
3. Evaluar la relevancia clínica del efecto observado.

Solución

- El valor p es menor que 0.05 , lo que indica significación estadística.
- El intervalo de confianza no incluye el valor nulo, lo que respalda la existencia de un efecto.
- La amplitud del intervalo sugiere una precisión moderada.
- La relevancia clínica depende de si una reducción de 8 mg/dL tiene impacto terapéutico en el contexto del paciente y las guías clínicas.

Interpretación clínica

Desde el punto de vista estadístico, existe evidencia de que el nuevo fármaco reduce el colesterol LDL en comparación con el tratamiento estándar. Sin embargo, la decisión clínica debe considerar si la magnitud del efecto justifica su uso, teniendo en cuenta beneficios adicionales, riesgos y costos asociados.

Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

CAPÍTULO V

Bioestadística Aplicada al Laboratorio Clínico

AUTOR: Chilán Santana Caleb Isaac



5 Bioestadística Aplicada al Laboratorio Clínico

5.1 Variabilidad biológica y analítica en pruebas de laboratorio

La correcta interpretación de los resultados de laboratorio clínico requiere una comprensión profunda de las fuentes de variabilidad que influyen en las mediciones biomédicas. En bioestadística aplicada al laboratorio clínico, la variabilidad se reconoce como un fenómeno inherente a los sistemas biológicos y a los procesos analíticos, y su adecuada cuantificación resulta esencial para diferenciar cambios fisiológicos reales de fluctuaciones debidas al proceso de medición. Ignorar estas fuentes de variación puede conducir a interpretaciones erróneas, diagnósticos incorrectos y decisiones clínicas inapropiadas (Westgard et al., 2010; Young, 2002).

Desde una perspectiva sanitaria, la variabilidad en los resultados de laboratorio se clasifica de manera general en dos grandes componentes: la variabilidad biológica y la variabilidad analítica. Ambos tipos interactúan entre sí y determinan la precisión, exactitud y utilidad clínica de las pruebas diagnósticas (Rifai et al., 2019).

5.1.1 Concepto de variabilidad en el laboratorio clínico

La variabilidad se define como el grado de dispersión de los valores obtenidos al medir repetidamente una misma magnitud bajo condiciones similares. En el laboratorio clínico, esta dispersión puede observarse tanto entre distintos individuos como en mediciones repetidas en un mismo paciente. La bioestadística proporciona herramientas para cuantificar esta variabilidad y evaluar su impacto en la interpretación clínica.

La variabilidad no debe entenderse únicamente como un error, sino como una característica natural de los sistemas biológicos y de los procedimientos de medición. El desafío del laboratorio clínico consiste en controlar y minimizar la variabilidad analítica, de modo que los cambios observados reflejen principalmente variaciones biológicas reales.

5.1.2 Variabilidad biológica

La variabilidad biológica corresponde a las fluctuaciones naturales de un analito dentro de un individuo y entre diferentes individuos. Estas variaciones

se deben a factores fisiológicos, genéticos, ambientales y temporales, y constituyen una fuente inevitable de dispersión en los resultados de laboratorio (Pineda-Tenor et al., 2013).

Se distinguen dos componentes principales de la variabilidad biológica:

- **Variabilidad intraindividual**, que refleja las fluctuaciones de un analito en un mismo individuo a lo largo del tiempo.
- **Variabilidad interindividual**, que representa las diferencias entre individuos de una población.

Ambos componentes son relevantes para establecer intervalos de referencia y para evaluar cambios clínicamente significativos en un paciente determinado.

Tabla 5-1. Componentes de la variabilidad biológica

Tipo de variabilidad	Descripción	Implicación clínica
Intraindividual	Fluctuaciones en un mismo individuo	Seguimiento del paciente
Interindividual	Diferencias entre personas	Intervalos de referencia

Nota. Componentes de la variabilidad biológica que influyen en la interpretación de resultados de laboratorio.

La variabilidad biológica intraindividual es especialmente importante en el monitoreo clínico, ya que un cambio estadísticamente pequeño puede ser clínicamente relevante si supera la variación habitual del paciente. En contraste, un valor fuera del intervalo de referencia poblacional no siempre indica patología si se mantiene dentro del rango habitual del individuo.

5.1.3 Variabilidad analítica

La variabilidad analítica se origina en el proceso de medición y está asociada al método analítico, al equipo, a los reactivos y al personal que realiza la prueba. A diferencia de la variabilidad biológica, este componente puede y

debe ser controlado mediante procedimientos estandarizados y sistemas de control de calidad.

Desde el punto de vista estadístico, la variabilidad analítica se expresa mediante medidas de dispersión como la desviación estándar y el coeficiente de variación, que permiten evaluar la precisión del método analítico.

El coeficiente de variación se define como:

$$CV = \frac{s}{\bar{x}} \times 100$$

donde s es la desviación estándar y $\bar{x}$ la media de las mediciones. Este indicador facilita la comparación de la variabilidad entre distintos analitos y métodos.

Tabla 5-2. Indicadores estadísticos de variabilidad analítica

Indicador	Definición	Uso en laboratorio
Desviación estándar	Dispersión absoluta	Precisión del método
Coefficiente de variación	Dispersión relativa	Comparación entre pruebas

Nota. Indicadores estadísticos utilizados para cuantificar la variabilidad analítica en el laboratorio clínico.

5.1.4 Relación entre variabilidad biológica y analítica

La utilidad clínica de una prueba de laboratorio depende de la relación entre la variabilidad analítica y la variabilidad biológica. Un principio ampliamente aceptado en bioestadística clínica establece que la variabilidad analítica debe ser considerablemente menor que la variabilidad biológica intraindividual, para que los cambios observados reflejen alteraciones fisiológicas reales y no fluctuaciones del método.

Cuando la variabilidad analítica es elevada, se dificulta la detección de cambios clínicamente relevantes y se reduce la confiabilidad del resultado. Por ello, los criterios de desempeño analítico se establecen considerando la magnitud de la variabilidad biológica del analito de interés.

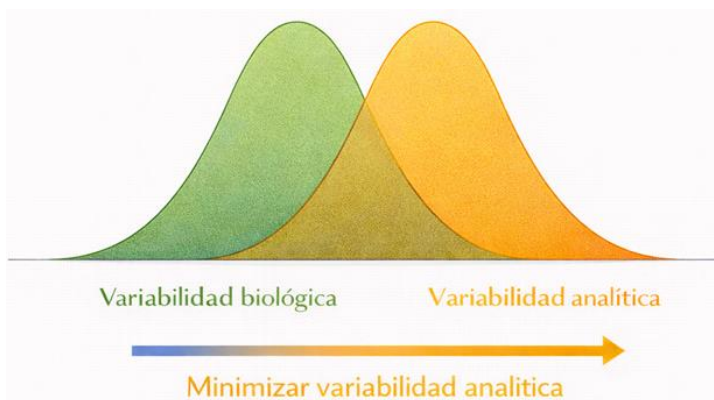


Figura 5-1. Relación entre variabilidad biológica y analítica

5.1.5 Implicaciones clínicas de la variabilidad

La comprensión de la variabilidad biológica y analítica tiene implicaciones directas en la práctica clínica. En el seguimiento de pacientes, por ejemplo, no basta con observar diferencias numéricas entre resultados sucesivos; es necesario determinar si dichas diferencias superan la variabilidad esperada y representan un cambio clínicamente significativo.

Asimismo, la variabilidad condiciona la interpretación de resultados cercanos a los puntos de decisión clínica, como valores límite o umbrales diagnósticos. Un resultado ligeramente superior o inferior a un punto de corte debe interpretarse considerando la imprecisión analítica y la variabilidad biológica del analito (Westgard, 2010).

5.1.6 Rol de la bioestadística en el control de la variabilidad

La bioestadística aplicada proporciona las herramientas conceptuales y matemáticas para cuantificar, monitorear y controlar la variabilidad en el laboratorio clínico. Mediante el análisis de datos repetidos, el cálculo de indicadores de dispersión y la comparación con criterios de desempeño, es posible evaluar la calidad analítica y asegurar la confiabilidad de los resultados reportados.

Este enfoque estadístico constituye la base para los sistemas de control de calidad interno y externo, que se desarrollarán en el siguiente epígrafe como estrategias fundamentales para garantizar la excelencia analítica en el laboratorio clínico.

5.2 Control de calidad interno y externo

El control de calidad constituye un componente esencial del funcionamiento del laboratorio clínico, ya que asegura que los resultados analíticos sean confiables, comparables y clínicamente útiles. Desde la bioestadística aplicada, el control de calidad se apoya en herramientas cuantitativas que permiten vigilar el desempeño analítico, detectar desviaciones sistemáticas o aleatorias y garantizar que la variabilidad analítica se mantenga dentro de límites aceptables. La correcta implementación del control de calidad interno y externo es indispensable para proteger la seguridad del paciente y sostener la validez diagnóstica de las pruebas de laboratorio.

En términos generales, el control de calidad en el laboratorio clínico se estructura en dos niveles complementarios. El control de calidad interno se orienta a la supervisión diaria del proceso analítico, mientras que el control de calidad externo evalúa el desempeño del laboratorio en comparación con otros laboratorios y con valores de referencia aceptados a nivel nacional o internacional.

5.2.1 Control de calidad interno

El control de calidad interno comprende el conjunto de procedimientos aplicados de forma rutinaria dentro del laboratorio para verificar que los métodos analíticos funcionen de manera estable y consistente a lo largo del tiempo. Su objetivo principal es detectar errores antes de que los resultados sean reportados al personal clínico, reduciendo así el riesgo de decisiones diagnósticas incorrectas.

Este tipo de control se basa en el análisis estadístico repetido de materiales de control con concentraciones conocidas. Los resultados obtenidos se comparan con valores esperados y se evalúan mediante indicadores de dispersión y reglas estadísticas que permiten identificar comportamientos anómalos del sistema analítico.

5.2.2 Materiales de control y parámetros estadísticos

Los materiales de control interno son muestras estables, similares a las muestras clínicas, que se analizan junto con los especímenes de los pacientes. A partir de mediciones repetidas, se calculan parámetros estadísticos fundamentales como la media, la desviación estándar y el coeficiente de variación, los cuales sirven de referencia para evaluar el desempeño del método.

El coeficiente de variación se expresa como:

$$CV = \frac{s}{\bar{x}} \times 100$$

donde *s* representa la desviación estándar y $\bar{x}$ la media de los resultados del control. Valores elevados del coeficiente de variación indican baja precisión analítica y requieren acciones correctivas.

Tabla 5-3. Parámetros estadísticos utilizados en el control de calidad interno

Parámetro	Significado	Utilidad
Media	Valor central del control	Referencia esperada
Desviación estándar	Dispersión absoluta	Evaluación de precisión
Coeficiente de variación	Dispersión relativa	Comparación entre métodos

Nota. Indicadores estadísticos empleados para vigilar la estabilidad y precisión de los métodos analíticos.

5.2.3 Gráficos de control

Una herramienta fundamental del control de calidad interno es el gráfico de control, en el cual los resultados del material de control se representan en función del tiempo. En estos gráficos se incluyen líneas que corresponden a la media y a límites de control definidos en múltiplos de la desviación estándar. Esta representación visual permite identificar tendencias, desplazamientos y aumentos de variabilidad que pueden indicar problemas analíticos incipientes.

Los límites de control más utilizados corresponden a ± 2 y ± 3 desviaciones estándar respecto a la media. Resultados que superan estos límites sugieren una

pérdida de control del proceso y requieren la suspensión del reporte de resultados hasta identificar y corregir la causa.

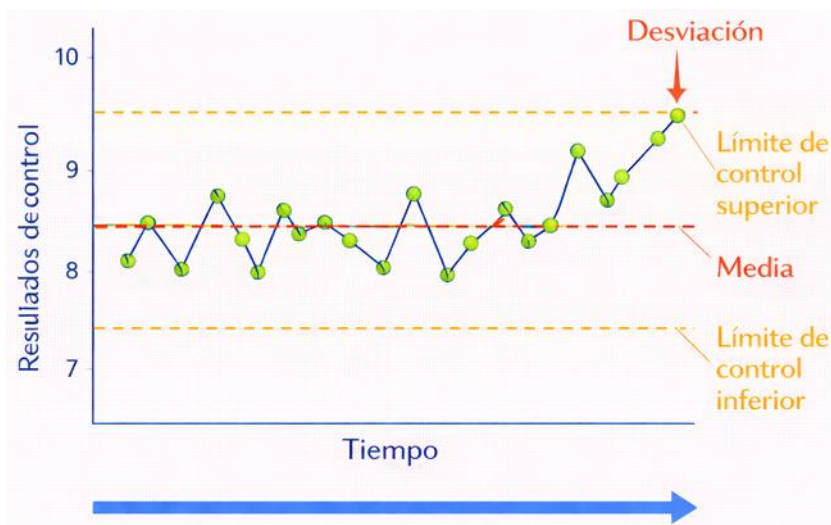


Figura 5-2. Gráfico de control para monitoreo analítico

5.2.4 Reglas estadísticas en el control de calidad interno

Para reforzar la detección temprana de errores, se utilizan reglas estadísticas que combinan información sobre la magnitud y la frecuencia de las desviaciones observadas. Estas reglas permiten distinguir entre errores aleatorios y errores sistemáticos, optimizando la sensibilidad del control de calidad.

Desde el punto de vista bioestadístico, estas reglas se basan en la probabilidad de que un resultado se desvíe de la media más allá de ciertos límites bajo condiciones normales de funcionamiento. La aplicación sistemática de estas reglas mejora la capacidad del laboratorio para mantener un desempeño analítico consistente.

5.2.5 Control de calidad externo

El control de calidad externo, también denominado evaluación externa del desempeño, consiste en la comparación periódica de los resultados del

laboratorio con los obtenidos por otros laboratorios que utilizan métodos similares. Este proceso permite evaluar la exactitud analítica y detectar sesgos sistemáticos que podrían no ser evidentes mediante el control interno (Rifai et al., 2019).

En este esquema, los laboratorios analizan muestras desconocidas y reportan los resultados a una entidad coordinadora, que posteriormente emite informes comparativos. La interpretación de estos informes requiere herramientas estadísticas que permitan evaluar la desviación respecto a valores asignados o consensuados (Fraser, 2001).

Tabla 5-4. Diferencias entre control de calidad interno y externo

Aspecto	Control interno	Control externo
Frecuencia	Diaria o rutinaria	Periódica
Objetivo	Precisión y estabilidad	Exactitud comparativa
Alcance	Proceso interno	Comparación interlaboratorial

Nota. Comparación funcional entre los sistemas de control de calidad interno y externo en el laboratorio clínico.

5.2.6 Interpretación estadística del control externo

En el control externo, el desempeño del laboratorio se evalúa mediante indicadores estadísticos que cuantifican la desviación respecto a un valor de referencia. Entre los más utilizados se encuentran las diferencias absolutas, las diferencias relativas y los puntajes estandarizados, que permiten juzgar si el resultado se encuentra dentro de límites aceptables.

Una interpretación adecuada de estos indicadores ayuda a identificar sesgos persistentes y a implementar acciones correctivas orientadas a mejorar la exactitud del método analítico.

5.2.7 Importancia clínica del control de calidad

El control de calidad interno y externo no solo tiene una función técnica, sino también una repercusión directa en la seguridad del paciente. Resultados analíticos incorrectos pueden conducir a diagnósticos erróneos, tratamientos

innecesarios o retrasos en la atención sanitaria. Por ello, la bioestadística aplicada al control de calidad constituye un soporte fundamental de la medicina de laboratorio moderna.

La integración de ambos sistemas de control garantiza que los resultados reportados sean consistentes a lo largo del tiempo y comparables entre diferentes laboratorios, fortaleciendo la confianza del personal clínico en la información analítica.

5.3 Intervalos de referencia y valores críticos

Los intervalos de referencia y los valores críticos constituyen herramientas fundamentales para la interpretación clínica de los resultados de laboratorio. Su correcta definición y uso permiten diferenciar resultados compatibles con la variación fisiológica normal de aquellos que indican una alteración significativa del estado de salud. Desde la bioestadística aplicada, ambos conceptos se sustentan en el análisis de la distribución de los datos, la estimación de percentiles y la consideración de la variabilidad biológica y analítica previamente descrita.

En la práctica clínica, un resultado analítico solo adquiere significado cuando se compara con un marco de referencia adecuado. Dicho marco debe construirse con criterios estadísticos sólidos y con una selección apropiada de la población de referencia, evitando interpretaciones simplistas o descontextualizadas.

5.3.1 Concepto de intervalo de referencia

El intervalo de referencia se define como el rango de valores que abarca una proporción determinada de resultados obtenidos en una población aparentemente sana. En el laboratorio clínico, se adopta de manera convencional el intervalo que comprende el 95 por ciento central de los valores, delimitado por los percentiles 2.5 y 97.5 de la distribución del analito.

Este enfoque reconoce que incluso en individuos sanos existe variabilidad biológica, y que valores extremos pueden presentarse sin implicar necesariamente patología. La bioestadística permite cuantificar esta variabilidad y establecer límites que reflejen la distribución real de los datos en la población de referencia.

5.3.2 Métodos estadísticos para el establecimiento de intervalos de referencia

Existen distintos métodos estadísticos para la determinación de intervalos de referencia, cuya elección depende de la forma de la distribución de los datos y del tamaño muestral. Cuando la distribución es aproximadamente normal, el intervalo de referencia puede estimarse mediante la media y la desviación estándar:

$$IR = \bar{x} \pm 1.96 s$$

donde $\bar{x}$ representa la media y s la desviación estándar de la muestra de referencia. Este método asume simetría y normalidad, condiciones que no siempre se cumplen en datos de laboratorio.

En situaciones donde la distribución no es normal, se recomienda el uso de métodos no paramétricos basados en percentiles, que no requieren supuestos estrictos sobre la forma de la distribución y ofrecen mayor robustez estadística.

Tabla 5-5. Métodos estadísticos para intervalos de referencia

Método	Supuesto principal	Aplicación
Paramétrico	Distribución normal	Analitos simétricos
No paramétrico	Sin supuestos de forma	Analitos asimétricos

Nota. Métodos estadísticos utilizados para estimar intervalos de referencia según la distribución de los datos.

5.3.3 Selección de la población de referencia

La validez de un intervalo de referencia depende de la adecuada selección de la población de referencia. Esta población debe representar individuos sin enfermedad relevante para el analito estudiado y considerar variables como edad, sexo y condiciones fisiológicas particulares. La bioestadística aplicada permite estratificar los datos y evaluar si es necesario establecer intervalos diferenciados para subgrupos específicos (Burtis et al., 2019).

Por ejemplo, parámetros hematológicos y bioquímicos pueden presentar diferencias significativas entre adultos y niños, o entre hombres y mujeres, lo que justifica intervalos de referencia específicos para cada grupo.

5.3.4 Valores críticos y su significado clínico

Los valores críticos se definen como resultados analíticos que indican una situación potencialmente grave y que requieren una acción clínica inmediata. A diferencia de los intervalos de referencia, los valores críticos no se establecen únicamente con criterios estadísticos, sino que incorporan juicio clínico y consenso profesional.

Desde el punto de vista bioestadístico, los valores críticos suelen situarse en los extremos de la distribución, más allá de los límites habituales del intervalo de referencia. Su identificación temprana y comunicación oportuna son esenciales para la seguridad del paciente y la efectividad de la atención sanitaria.

Tabla 5-6. Diferencias entre intervalos de referencia y valores críticos

Característica	Intervalos de referencia	Valores críticos
Base principal	Distribución poblacional	Riesgo clínico
Uso	Interpretación general	Acción inmediata
Definición	Estadística	Estadística y clínica

Nota. Comparación conceptual entre intervalos de referencia y valores críticos en el laboratorio clínico.

5.3.5 Influencia de la variabilidad en la interpretación

La interpretación de resultados cercanos a los límites del intervalo de referencia debe considerar la variabilidad biológica y analítica. Un valor ligeramente fuera del intervalo no siempre implica patología, especialmente si la variación observada se encuentra dentro del rango esperado para el individuo. La bioestadística aporta criterios para evaluar si un cambio es estadísticamente significativo o clínicamente relevante (Fraser, 2001).

En el caso de valores críticos, incluso pequeñas desviaciones pueden tener consecuencias clínicas importantes, por lo que la tolerancia a la variabilidad es menor y se requieren protocolos estrictos de verificación y comunicación.

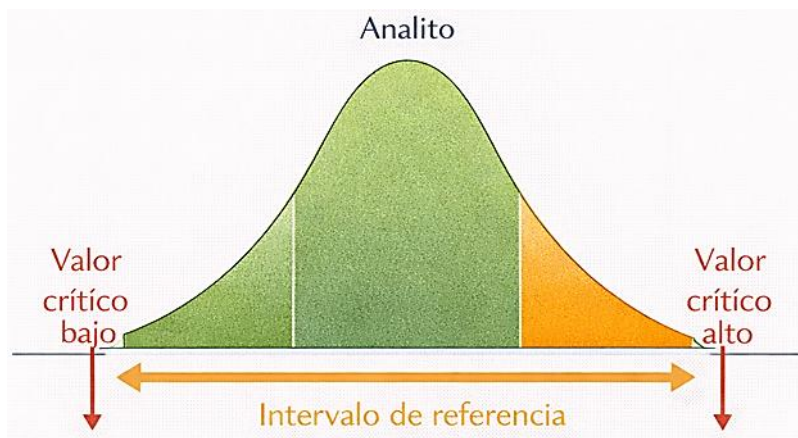


Figura 5-3. Intervalos de referencia y valores críticos

Propuesta de figura que muestre una distribución de un analito con el intervalo de referencia central y los valores críticos ubicados en los extremos de la curva.

5.3.6 Importancia clínica y bioestadística

La correcta utilización de intervalos de referencia y valores críticos fortalece la capacidad del laboratorio clínico para apoyar decisiones médicas fundamentadas. La bioestadística aplicada permite construir intervalos confiables, evaluar su validez y comprender sus limitaciones, contribuyendo a una interpretación más precisa y segura de los resultados.

5.4 Validación estadística de métodos analíticos

La validación estadística de métodos analíticos es un proceso fundamental en el laboratorio clínico, orientado a demostrar que un método de medición es adecuado para su propósito clínico. Desde la bioestadística aplicada, la validación permite evaluar de manera objetiva el desempeño analítico, garantizando que los resultados obtenidos sean confiables, reproducibles y clínicamente interpretables. Un método no validado puede generar resultados

imprecisos o sesgados, comprometiendo la calidad diagnóstica y la seguridad del paciente.

En el contexto sanitario, la validación estadística no se limita a un requisito técnico, sino que constituye una responsabilidad ética y profesional del laboratorio clínico. Este proceso se apoya en el análisis cuantitativo de datos experimentales y en criterios estadísticos que permiten juzgar si un método cumple con estándares aceptables de desempeño.

5.4.1 Objetivo de la validación de métodos analíticos

El objetivo principal de la validación es confirmar que un método analítico produce resultados exactos y precisos dentro de un rango clínicamente relevante. Desde el punto de vista bioestadístico, la validación busca cuantificar y controlar las fuentes de error analítico, asegurando que la variabilidad asociada al método no interfiera con la detección de cambios biológicos reales.

La validación se aplica tanto a métodos nuevos como a métodos previamente establecidos que han sufrido modificaciones en equipos, reactivos o procedimientos. En todos los casos, el análisis estadístico es el soporte esencial para la toma de decisiones sobre la aceptación del método.

5.4.2 Parámetros estadísticos evaluados en la validación

La validación estadística de un método analítico se basa en la evaluación sistemática de varios parámetros fundamentales. Entre los más relevantes se encuentran la precisión, la exactitud, la linealidad y el rango de medición. Cada uno de estos parámetros se cuantifica mediante indicadores estadísticos específicos que permiten evaluar el desempeño del método de manera objetiva.

La precisión se relaciona con la dispersión de los resultados obtenidos en mediciones repetidas, mientras que la exactitud expresa la cercanía entre el valor medido y el valor verdadero o de referencia. Ambos conceptos son complementarios y deben analizarse de forma conjunta para una validación completa.

Tabla 5-7. Parámetros estadísticos en la validación de métodos analíticos

Parámetro	Descripción	Enfoque estadístico
Precisión	Repetibilidad de resultados	Desviación estándar, CV
Exactitud	Cercanía al valor verdadero	Sesgo, diferencia media
Linealidad	Proporcionalidad respuesta-concentración	Regresión lineal
Rango	Intervalo de medición válido	Evaluación experimental

Nota. Principales parámetros estadísticos evaluados durante la validación de métodos analíticos en el laboratorio clínico.

5.4.3 Evaluación de la precisión

La precisión se evalúa mediante el análisis de mediciones repetidas de una misma muestra bajo condiciones controladas. Desde la bioestadística, la precisión se cuantifica utilizando la desviación estándar y el coeficiente de variación, que permiten estimar la dispersión de los resultados.

Un método con baja dispersión presenta alta precisión, lo que indica estabilidad analítica. Sin embargo, una alta precisión no garantiza necesariamente exactitud, por lo que ambos aspectos deben evaluarse de manera independiente.

5.4.4 Evaluación de la exactitud y el sesgo

La exactitud se analiza comparando los resultados del método evaluado con un valor de referencia aceptado. La diferencia sistemática entre ambos valores se denomina sesgo. Desde el punto de vista estadístico, el sesgo se calcula como la diferencia media entre los resultados observados y el valor de referencia (Fraser, 2001).

Un sesgo elevado indica una desviación sistemática que puede afectar la interpretación clínica de los resultados, especialmente cuando los valores se aproximan a puntos de decisión diagnóstica.

5.4.5 Análisis de linealidad

La linealidad evalúa la capacidad del método para producir resultados proporcionales a la concentración del analito dentro de un rango determinado. Este análisis se realiza mediante técnicas de regresión lineal, donde se examina la relación entre la concentración conocida y la respuesta medida (Motulsky, 2015).

La ecuación general de la regresión lineal utilizada en este contexto es:

$$y = \beta_0 + \beta_1 x$$

donde y representa la respuesta analítica y x la concentración del analito. Una relación lineal adecuada respalda la validez del método dentro del rango evaluado.

5.4.6 Criterios de aceptación y decisión estadística

La aceptación de un método analítico se basa en la comparación de los parámetros estadísticos obtenidos con criterios previamente establecidos. Estos criterios pueden definirse a partir de la variabilidad biológica del analito o de estándares de desempeño reconocidos en medicina de laboratorio.

Desde la bioestadística aplicada, esta comparación permite decidir si el método es apto para uso clínico, requiere ajustes o debe ser rechazado. La decisión final debe considerar no solo los resultados numéricos, sino también su impacto potencial en la interpretación clínica.



Figura 5-4. Proceso de validación estadística de un método analítico

5.4.7 Importancia clínica de la validación estadística

La validación estadística de métodos analíticos garantiza que los resultados reportados por el laboratorio clínico sean confiables y comparables en el

tiempo. Desde el punto de vista sanitario, este proceso protege al paciente frente a errores diagnósticos y fortalece la confianza del personal clínico en la información analítica.

5.5 Concordancia diagnóstica: sensibilidad, especificidad y valores predictivos

La evaluación del desempeño diagnóstico de las pruebas de laboratorio es un componente central de la bioestadística aplicada a las ciencias de la salud. Más allá de la precisión analítica del método, resulta indispensable determinar en qué medida una prueba permite identificar correctamente a los individuos con y sin una determinada condición clínica. En este contexto, la concordancia diagnóstica se fundamenta en indicadores estadísticos que cuantifican la capacidad discriminativa de una prueba, entre los cuales destacan la sensibilidad, la especificidad y los valores predictivos.

Estos indicadores permiten evaluar la utilidad clínica real de una prueba diagnóstica y orientan su uso en distintos escenarios sanitarios, como tamizaje, confirmación diagnóstica o seguimiento de enfermedades.

5.5.1 Concepto de concordancia diagnóstica

La concordancia diagnóstica se refiere al grado de acuerdo entre los resultados de una prueba diagnóstica y el estado real de la enfermedad, determinado mediante un criterio de referencia aceptado. Desde la bioestadística aplicada, este acuerdo se cuantifica comparando los resultados de la prueba con la condición verdadera del individuo, lo que permite clasificar los resultados en verdaderos positivos, verdaderos negativos, falsos positivos y falsos negativos.

Este enfoque reconoce que ninguna prueba diagnóstica es perfecta y que la evaluación cuantitativa de sus aciertos y errores es esencial para interpretar sus resultados en la práctica clínica.

5.5.2 Tabla de contingencia diagnóstica

La base estadística para el cálculo de los indicadores de concordancia diagnóstica es la tabla de contingencia 2×2 , que resume la relación entre el resultado de la prueba y la condición real del individuo.

Tabla 5-8. Tabla de contingencia para evaluación diagnóstica

Resultado de la prueba	Enfermedad presente	Enfermedad ausente
Positivo	Verdadero positivo	Falso positivo
Negativo	Falso negativo	Verdadero negativo

Nota. Estructura básica para cuantificar la concordancia entre resultados diagnósticos y el estado real de la enfermedad.

5.5.3 Sensibilidad

La sensibilidad se define como la capacidad de una prueba para identificar correctamente a los individuos que realmente presentan la enfermedad. Desde el punto de vista estadístico, representa la proporción de verdaderos positivos entre todos los individuos con la condición clínica (Gordis, 2014).

La expresión matemática de la sensibilidad es:

$$\text{Sensibilidad} = \frac{VP}{VP + FN}$$

donde VP corresponde a verdaderos positivos y FN a falsos negativos.

Una prueba con alta sensibilidad es especialmente útil en contextos de tamizaje, ya que reduce la probabilidad de clasificar erróneamente como sanos a individuos enfermos.

5.5.4 Especificidad

La especificidad expresa la capacidad de una prueba para identificar correctamente a los individuos que no presentan la enfermedad. Estadísticamente, se define como la proporción de verdaderos negativos entre todos los individuos sin la condición clínica (Altman, 1991).

Su expresión matemática es:

$$\text{Especificidad} = \frac{VN}{VN + FP}$$

donde VN representa verdaderos negativos y FP falsos positivos.

Las pruebas con alta especificidad son particularmente valiosas para confirmar diagnósticos, ya que minimizan la probabilidad de clasificar erróneamente como enfermos a individuos sanos.

5.5.5 Valores predictivos

A diferencia de la sensibilidad y la especificidad, que dependen fundamentalmente de las características intrínsecas de la prueba, los valores predictivos incorporan la prevalencia de la enfermedad en la población evaluada. Estos indicadores responden a la pregunta clínica de cuán probable es que un resultado positivo o negativo refleje el estado real del paciente (Gordis, 2014).

El valor predictivo positivo se define como:

$$VPP = \frac{VP}{VP + FP}$$

El valor predictivo negativo se expresa como:

$$VPN = \frac{VN}{VN + FN}$$

Estos indicadores son especialmente relevantes en la práctica clínica diaria, ya que influyen directamente en la interpretación de los resultados reportados por el laboratorio.

Tabla 5-9. Indicadores de concordancia diagnóstica

Indicador	Definición	Utilidad clínica
Sensibilidad	Identifica enfermos	Tamizaje
Especificidad	Identifica sanos	Confirmación
VPP	Probabilidad de enfermedad	Interpretación clínica
VPN	Probabilidad de ausencia	Seguimiento

Nota. Principales indicadores estadísticos utilizados para evaluar el desempeño diagnóstico de pruebas de laboratorio.

5.5.6 Influencia de la prevalencia en los valores predictivos

La prevalencia de la enfermedad en la población evaluada tiene un impacto directo sobre los valores predictivos. En poblaciones con baja prevalencia, incluso pruebas con alta sensibilidad y especificidad pueden presentar valores predictivos positivos bajos. Este fenómeno debe considerarse cuidadosamente al interpretar resultados de pruebas diagnósticas en distintos contextos sanitarios (Gordis, 2014).

Desde la bioestadística aplicada, esta relación explica por qué una misma prueba puede tener desempeños clínicos distintos según el escenario epidemiológico en el que se utilice.

5.5.7 Relevancia clínica de la concordancia diagnóstica

La evaluación de la concordancia diagnóstica permite seleccionar pruebas apropiadas para cada propósito clínico y optimizar los algoritmos diagnósticos. Comprender las fortalezas y limitaciones de la sensibilidad, la especificidad y los valores predictivos contribuye a una interpretación más precisa de los resultados y a una mejor toma de decisiones en salud (Rosner, 2016).

5.6 Curvas ROC y análisis de desempeño de pruebas diagnósticas

El análisis de curvas ROC constituye una de las herramientas más robustas de la bioestadística aplicada para evaluar el desempeño global de las pruebas diagnósticas. A diferencia de los indicadores clásicos como sensibilidad y especificidad, que se calculan para un único punto de decisión, las curvas ROC permiten analizar el comportamiento de una prueba a lo largo de todos los posibles puntos de corte. Esta perspectiva integral resulta especialmente valiosa en el laboratorio clínico, donde muchos analitos se expresan como variables continuas y la elección del punto de decisión tiene implicaciones clínicas directas.

Desde el punto de vista sanitario, las curvas ROC facilitan la comparación objetiva entre métodos diagnósticos, apoyan la selección de puntos de corte clínicamente adecuados y contribuyen a optimizar la toma de decisiones médicas basadas en evidencia cuantitativa.

5.6.1 Fundamento estadístico de la curva ROC

La curva ROC se construye representando gráficamente la sensibilidad frente a uno menos la especificidad para distintos puntos de corte de una prueba diagnóstica. Cada punto de la curva corresponde a un umbral específico utilizado para clasificar los resultados como positivos o negativos. Este enfoque permite visualizar el compromiso existente entre la capacidad de detectar verdaderos positivos y la probabilidad de generar falsos positivos.

Desde la bioestadística, la curva ROC se interpreta como una representación de la capacidad discriminativa de una prueba para distinguir entre individuos con y sin la condición clínica de interés.

5.6.2 Área bajo la curva ROC

El área bajo la curva ROC, conocida como AUC, es un indicador resumen del desempeño diagnóstico global de una prueba. Este valor varía entre 0.5 y 1.0, donde un AUC de 0.5 indica ausencia de capacidad discriminativa, equivalente a una clasificación aleatoria, mientras que un AUC de 1.0 representa una discriminación perfecta.

Desde el punto de vista clínico, el AUC permite comparar pruebas diagnósticas de manera objetiva, independientemente del punto de corte elegido. Cuanto mayor es el AUC, mayor es la capacidad de la prueba para diferenciar correctamente entre estados de salud y enfermedad.

Tabla 5-10. Interpretación general del área bajo la curva ROC

AUC	Interpretación
0.50	Sin discriminación
0.60 – 0.70	Discriminación baja
0.70 – 0.80	Discriminación aceptable
0.80 – 0.90	Discriminación buena
> 0.90	Discriminación excelente

Nota. Guía orientativa para interpretar la capacidad discriminativa de una prueba diagnóstica mediante el área bajo la curva ROC.

5.6.3 Selección del punto de corte óptimo

La elección del punto de corte es una decisión crítica en el uso clínico de pruebas diagnósticas. Desde la bioestadística aplicada, el punto de corte óptimo se selecciona considerando el equilibrio entre sensibilidad y especificidad, así como las consecuencias clínicas de los errores diagnósticos. Un punto de corte bajo incrementa la sensibilidad, pero reduce la especificidad, mientras que un punto de corte alto produce el efecto contrario.

La selección del umbral adecuado debe basarse en el contexto clínico, el propósito de la prueba y la gravedad de los errores de clasificación. En programas de tamizaje se prioriza la sensibilidad, mientras que en pruebas confirmatorias se busca maximizar la especificidad.

5.6.4 Comparación de pruebas diagnósticas mediante curvas ROC

Las curvas ROC permiten comparar el desempeño de dos o más pruebas diagnósticas aplicadas al mismo problema clínico. Desde una perspectiva estadística, la comparación se realiza evaluando las diferencias entre las áreas bajo las curvas correspondientes. Una prueba con mayor AUC presenta una capacidad discriminativa superior.

Este enfoque es ampliamente utilizado en el laboratorio clínico para seleccionar métodos analíticos, evaluar nuevas tecnologías diagnósticas y validar procedimientos alternativos frente a métodos de referencia.

Tabla 5-11. Uso de curvas ROC en evaluación diagnóstica

Aplicación	Utilidad
Comparación de métodos	Selección del mejor desempeño
Definición de puntos de corte	Optimización clínica
Validación diagnóstica	Evaluación global

Nota. Principales aplicaciones de las curvas ROC en la evaluación del desempeño de pruebas diagnósticas.

5.6.5 Relación entre curvas ROC y prevalencia

Una ventaja importante del análisis ROC es que el AUC no depende de la prevalencia de la enfermedad, a diferencia de los valores predictivos. Esto permite evaluar la capacidad discriminativa intrínseca de la prueba sin que el resultado se vea influido por la frecuencia de la condición en la población estudiada (Pepe, 2003).

No obstante, aunque el AUC sea independiente de la prevalencia, la interpretación clínica del punto de corte seleccionado debe considerar siempre el contexto epidemiológico y la población a la que se aplica la prueba.

5.6.6 Limitaciones del análisis ROC

A pesar de su utilidad, el análisis ROC presenta limitaciones que deben reconocerse. La curva ROC no proporciona información directa sobre los costos clínicos asociados a los errores diagnósticos ni sobre la prevalencia de la enfermedad. Además, dos pruebas con AUC similares pueden diferir en su desempeño clínico en regiones específicas de la curva, lo que requiere un análisis detallado más allá del valor global del AUC (Pepe, 2003).

Desde la bioestadística aplicada, estas limitaciones subrayan la necesidad de interpretar las curvas ROC como parte de un análisis integral del desempeño diagnóstico, complementado con otros indicadores.

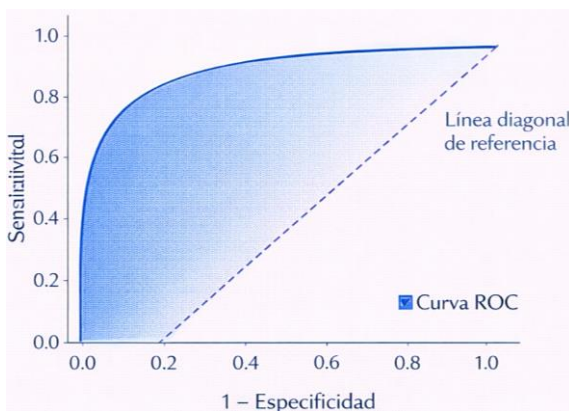


Figura 5-5. Curva ROC en evaluación de una prueba diagnóstica

5.6.7 Importancia clínica del análisis ROC

El análisis de curvas ROC aporta una visión integral del desempeño diagnóstico de las pruebas de laboratorio clínico. Al permitir evaluar simultáneamente sensibilidad y especificidad para múltiples puntos de corte, esta herramienta fortalece la selección y validación de pruebas diagnósticas, contribuyendo a una práctica clínica más segura y eficiente (Zweig y Campbell, 1993; Pepe, 2003).

En conjunto con los indicadores clásicos de concordancia diagnóstica, las curvas ROC permiten integrar la evaluación estadística con las necesidades clínicas, mejorando la calidad del diagnóstico y la toma de decisiones sanitarias.

5.7 Aplicaciones prácticas en bioquímica, hematología y microbiología

La bioestadística aplicada al laboratorio clínico alcanza su mayor utilidad cuando los principios analíticos y los indicadores de desempeño se integran en escenarios reales de bioquímica, hematología y microbiología. En estas áreas, la correcta interpretación de resultados depende de la comprensión de la variabilidad, del control de calidad, de los intervalos de referencia y del desempeño diagnóstico de las pruebas. La aplicación coherente de métodos estadísticos permite transformar datos analíticos en información clínica confiable, comparable y útil para la toma de decisiones sanitarias.

Este epígrafe presenta aplicaciones prácticas que ilustran cómo los conceptos desarrollados a lo largo del capítulo se emplean en el laboratorio clínico, culminando con un ejercicio integral que articula control de calidad, concordancia diagnóstica y análisis ROC.

5.7.1 Aplicaciones en bioquímica clínica

En bioquímica clínica, los analitos se expresan mayoritariamente como variables continuas, lo que exige un manejo cuidadoso de la variabilidad analítica y biológica. La bioestadística aplicada se utiliza para establecer intervalos de referencia, evaluar la precisión del método y definir puntos de decisión clínica para analitos como glucosa, creatinina, colesterol y enzimas hepáticas.

La selección de intervalos de referencia adecuados y la evaluación del sesgo analítico resultan esenciales para evitar errores diagnósticos, especialmente cuando los resultados se sitúan cerca de valores límite. Asimismo, el uso de indicadores estadísticos permite comparar métodos alternativos y garantizar la trazabilidad de los resultados a lo largo del tiempo.

Tabla 5-12. Uso de la bioestadística en bioquímica clínica

Aplicación	Indicador estadístico	Finalidad
Precisión analítica	Desviación estándar, CV	Estabilidad del método
Intervalos de referencia	Percentiles	Interpretación clínica
Puntos de decisión	Curvas ROC	Optimización diagnóstica

Nota. Principales aplicaciones de herramientas bioestadísticas en el análisis de pruebas bioquímicas.

5.7.2 Aplicaciones en hematología

En hematología, la bioestadística aplicada apoya la interpretación de parámetros como hemoglobina, hematocrito, recuentos celulares y volúmenes corpusculares. Estos analitos presentan variabilidad fisiológica relevante según edad, sexo y condiciones clínicas, lo que exige intervalos de referencia específicos y control analítico riguroso.

La evaluación estadística de la precisión es especialmente importante en hematología automatizada, donde pequeñas desviaciones pueden afectar la clasificación de anemias o trastornos hematológicos. Además, los estudios de concordancia diagnóstica permiten evaluar el desempeño de analizadores automatizados frente a métodos de referencia.

Tabla 5-13. Enfoque bioestadístico en hematología

Parámetro	Enfoque estadístico	Impacto clínico
Recuentos celulares	Precisión y exactitud	Clasificación diagnóstica
Índices eritrocitarios	Intervalos de referencia	Tipificación de anemias
Comparación de métodos	Concordancia diagnóstica	Validación instrumental

Nota. Integración de indicadores estadísticos en la interpretación de pruebas hematológicas.

5.7.3 Aplicaciones en microbiología clínica

En microbiología clínica, muchas pruebas producen resultados categóricos, como positivo o negativo, o semicuantitativos, como títulos o recuentos. En este contexto, la bioestadística aplicada se centra en la evaluación del desempeño diagnóstico mediante sensibilidad, especificidad y valores predictivos, así como en el análisis de curvas ROC cuando se emplean umbrales cuantitativos (Gordis, 2014).

La interpretación estadística es fundamental para validar pruebas rápidas, métodos moleculares y sistemas automatizados, especialmente cuando se introducen nuevas tecnologías diagnósticas. La selección del punto de corte adecuado tiene implicaciones directas en la detección temprana de infecciones y en el control epidemiológico.

Tabla 5-14. Aplicación de indicadores diagnósticos en microbiología

Indicador	Uso principal	Escenario clínico
Sensibilidad	Detección temprana	Tamizaje
Especificidad	Confirmación	Diagnóstico definitivo
Curvas ROC	Selección de umbral	Métodos cuantitativos

Nota. Uso de indicadores de desempeño diagnóstico en pruebas microbiológicas.

5.7.4 Integración bioestadística en la práctica del laboratorio clínico

La integración de herramientas bioestadísticas en bioquímica, hematología y microbiología fortalece la calidad analítica y la seguridad del paciente. El análisis sistemático de la variabilidad, la validación de métodos y la evaluación del desempeño diagnóstico permiten al laboratorio clínico cumplir su rol como soporte esencial de la medicina basada en evidencia.

Esta integración requiere una formación sólida del personal en interpretación estadística y un enfoque sistemático que vincule los resultados analíticos con su impacto clínico.

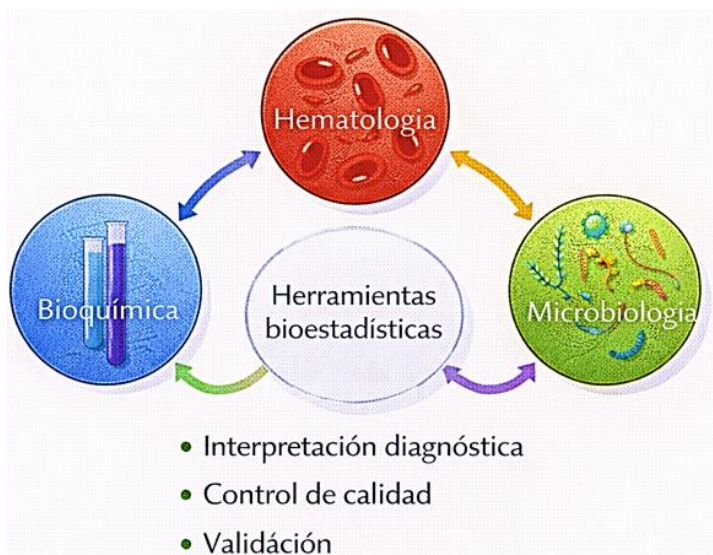


Figura 5-6. Integración de la bioestadística en áreas del laboratorio clínico

Ejercicio

Un laboratorio clínico hospitalario implementa un nuevo método automatizado para la determinación de glucosa sérica con fines diagnósticos y de seguimiento de pacientes con riesgo metabólico. Antes de su uso rutinario, el laboratorio debe evaluar estadísticamente el método, su desempeño analítico y su utilidad diagnóstica, aplicando principios de bioestadística clínica.

El estudio incluye:

- Evaluación de variabilidad biológica y analítica
- Control de calidad interno
- Definición de intervalos de referencia
- Validación estadística del método
- Evaluación de concordancia diagnóstica
- Análisis mediante curvas ROC

5.7.5 Datos disponibles

Mediciones repetidas de un control interno (glucosa mg/dL)

Medición	Valor
1	98
2	101
3	99
4	100
5	102
6	97
7	100
8	99
9	101
10	98

Valor asignado del control: **100 mg/dL**

Resultados en población de referencia sana (n = 20)

Paciente	Glucosa	9	93
1	85	10	94
2	90	11	86
3	88	12	90
4	92	13	88
5	95	14	92
6	87	15	89
7	91	16	91
8	89	17	87

18	93	20	88
19	90		

Comparación diagnóstica con criterio clínico (n = 120)

Resultado prueba	Enfermedad presente	Enfermedad ausente
Positivo	36	12
Negativo	9	63

Actividad 1: Variabilidad analítica

Cálculo de parámetros estadísticos

- Media:

$$\bar{x} = \frac{995}{10} = 99.5 \text{ mg/dL}$$

- Desviación estándar:

$$s = 1.71 \text{ mg/dL}$$

- Coefficiente de variación:

$$CV = \frac{1.71}{99.5} \times 100 = 1.72\%$$

Interpretación

El coeficiente de variación es bajo, lo que indica alta precisión analítica. La variabilidad analítica es inferior a la variabilidad biológica esperada para glucosa, cumpliendo criterios de aceptabilidad clínica.

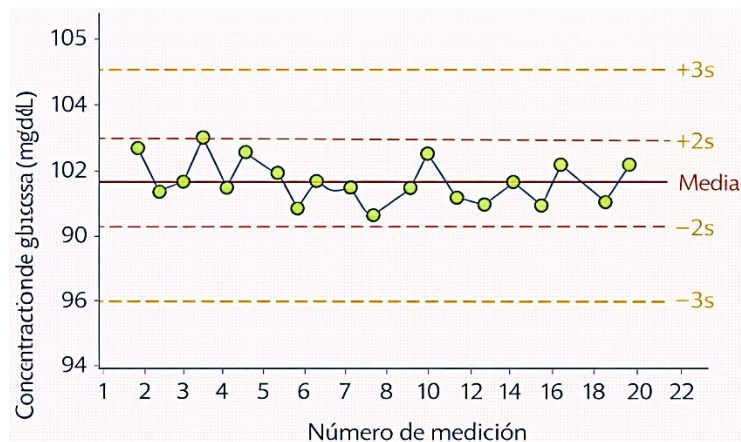
Actividad 2: Control de calidad interno

Se establecen límites de control:

- Media = 99.5 mg/dL

- $\pm 2s$: 96.1 – 102.9 mg/dL
- $\pm 3s$: 94.4 – 104.6 mg/dL

Gráfico de control de Levey-Jennings:



Interpretación

Todos los valores se encuentran dentro de $\pm 2s$, sin tendencias ni desplazamientos.

El método se considera bajo control estadístico.

Actividad 3: Intervalo de referencia

Cálculo estadístico

- Media poblacional: 90 mg/dL
- Desviación estándar: 3 mg/dL

Intervalo de referencia (95%):

$$IR = 90 \pm 1.96(3) = [84.1; 95.9] \text{ mg/dL}$$

Interpretación clínica

Valores dentro de este intervalo se consideran compatibles con normoglucemia en población adulta sana.

Actividad 4: Validación del método analítico

Evaluación de exactitud (sesgo)

$$\text{Sesgo} = 99.5 - 100 = -0.5 \text{ mg/dL}$$

Interpretación

El sesgo es mínimo y clínicamente irrelevante.

El método presenta buena exactitud.

Actividad 5: Concordancia diagnóstica

Cálculo de indicadores

- Sensibilidad:

$$\frac{36}{36 + 9} = 0.80$$

- Especificidad:

$$\frac{63}{63 + 12} = 0.84$$

- Valor predictivo positivo:

$$\frac{36}{48} = 0.75$$

- Valor predictivo negativo:

$$\frac{63}{72} = 0.88$$

Tabla resumen

Indicador	Valor
Sensibilidad	80%

Especificidad 84%

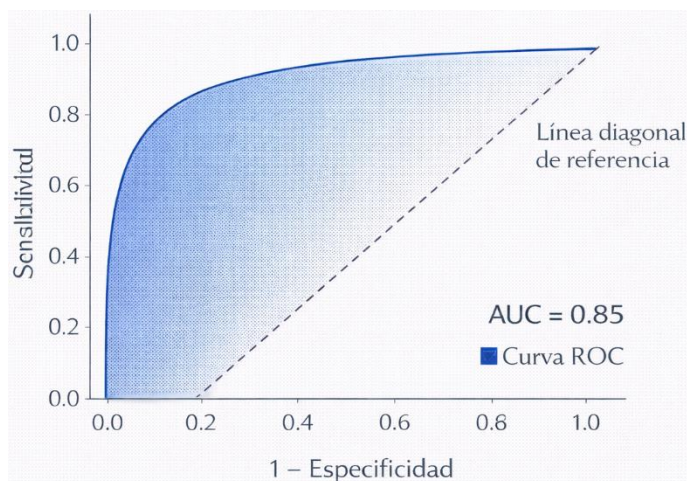
VPP 75%

VPN 88%

Actividad 6: Curva ROC

Análisis conceptual

- Se evalúan múltiples puntos de corte de glucosa
- Se construye la curva Sensibilidad vs (1 – Especificidad)
- Área bajo la curva estimada: **AUC = 0.85**



Interpretación

El método presenta buena capacidad discriminativa.

El punto de corte puede ajustarse según prioridad clínica:

- Tamizaje: mayor sensibilidad
- Confirmación: mayor especificidad

Conclusión del ejercicio

El análisis bioestadístico completo demuestra que el método de glucosa:

- Presenta variabilidad analítica controlada
- Cumple criterios de control de calidad interno
- Permite establecer intervalos de referencia confiables
- Está estadísticamente validado
- Muestra adecuado desempeño diagnóstico
- Posee buena capacidad discriminativa global.

Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

CAPÍTULO VI

Análisis Epidemiológico y Estudios Clínicos

AUTOR: Mina Ortiz Jhon Brian



6 Análisis Epidemiológico y Estudios Clínicos

6.1 Tipos de estudios epidemiológicos: descriptivos, analíticos y experimentales

El análisis epidemiológico constituye un pilar fundamental de la investigación sanitaria, ya que permite estudiar la distribución, los determinantes y las consecuencias de los eventos relacionados con la salud en poblaciones humanas. La epidemiología moderna se apoya en diseños de estudio bien definidos, cuya correcta selección condiciona la validez de los resultados, la interpretación estadística y la utilidad de las conclusiones para la práctica clínica y la salud pública. Desde la bioestadística aplicada, los estudios epidemiológicos se clasifican en tres grandes categorías: descriptivos, analíticos y experimentales, cada una con objetivos, alcances y niveles de evidencia diferenciados (Gordis, 2014; Rothman et al., 2008).

La comprensión de estos tipos de estudios resulta esencial para profesionales de las ciencias de la salud, ya que orienta tanto la lectura crítica de publicaciones científicas como el diseño de investigaciones propias, asegurando coherencia metodológica y rigor analítico.

6.1.1 Estudios epidemiológicos descriptivos

Los estudios descriptivos tienen como objetivo principal caracterizar la distribución de eventos de salud en una población, sin establecer relaciones causales. Estos estudios responden a preguntas básicas como quién se enferma, dónde ocurre el evento y cuándo se presenta, proporcionando una visión inicial del problema sanitario (Gordis, 2014).

Desde el punto de vista estadístico, los estudios descriptivos se apoyan principalmente en medidas de frecuencia, tasas y proporciones, así como en representaciones gráficas que permiten identificar patrones espaciales, temporales y poblacionales. Su utilidad radica en la generación de hipótesis que posteriormente pueden ser evaluadas mediante diseños analíticos más complejos. Entre los principales tipos de estudios descriptivos se incluyen los reportes de casos, las series de casos y los estudios transversales. Aunque estos

diseños no permiten establecer relaciones de causalidad, son esenciales para la vigilancia epidemiológica y la detección temprana de problemas emergentes de salud.

Tabla 6-1. Características de los estudios epidemiológicos descriptivos

Tipo de estudio	Característica principal	Uso en salud
Reporte de casos	Descripción individual	Eventos inusuales
Serie de casos	Grupo sin comparación	Nuevas enfermedades
Estudio transversal	Medición en un momento	Situación de salud

Nota. Principales diseños descriptivos utilizados para caracterizar eventos de salud en poblaciones humanas.

Los estudios transversales, en particular, ocupan un lugar destacado en epidemiología, ya que permiten estimar la prevalencia de enfermedades y factores de riesgo en poblaciones definidas. Desde la bioestadística aplicada, estos estudios requieren un diseño muestral adecuado y un análisis cuidadoso de posibles sesgos de selección y de información.

6.1.2 Estudios epidemiológicos analíticos

Los estudios analíticos buscan evaluar asociaciones entre exposiciones y desenlaces de salud, con el objetivo de identificar posibles determinantes de enfermedad. A diferencia de los estudios descriptivos, estos diseños incorporan grupos de comparación y permiten cuantificar la fuerza de la asociación mediante medidas estadísticas específicas (Rothman et al., 2008).

Los principales diseños analíticos son los estudios de cohortes y los estudios de casos y controles. Ambos se apoyan en el análisis comparativo y en el cálculo de medidas de asociación, como el riesgo relativo y el odds ratio. Desde el punto de vista bioestadístico, los estudios analíticos requieren un control riguroso de sesgos y factores de confusión, así como una planificación adecuada del tamaño muestral para garantizar potencia estadística suficiente.

Tabla 6-2. Diseños epidemiológicos analíticos

Diseño	Dirección del estudio	Medida típica
Cohorte	Exposición → evento	Riesgo relativo
Casos y controles	Evento → exposición	Odds ratio

Nota. Diseños analíticos más utilizados para evaluar asociaciones entre exposiciones y eventos de salud.

En los estudios de cohortes, los individuos se clasifican según su estado de exposición y se siguen en el tiempo para observar la ocurrencia del evento de interés. Este diseño permite estimar directamente la incidencia y evaluar relaciones temporales, lo que fortalece la inferencia epidemiológica.

Por su parte, los estudios de casos y controles comparan la frecuencia de exposición entre individuos con la enfermedad y un grupo de referencia sin la enfermedad. Este diseño resulta especialmente útil para el estudio de enfermedades poco frecuentes o con largos periodos de latencia.

6.1.3 Estudios epidemiológicos experimentales

Los estudios experimentales representan el nivel más alto de evidencia en epidemiología clínica, ya que el investigador interviene activamente asignando la exposición o tratamiento a los participantes. El diseño experimental más representativo es el ensayo clínico controlado, ampliamente utilizado para evaluar la eficacia y seguridad de intervenciones terapéuticas o preventivas.

Desde la bioestadística aplicada, los estudios experimentales se caracterizan por la asignación aleatoria, el uso de grupos control y la aplicación de análisis inferenciales robustos que permiten minimizar sesgos y maximizar la validez interna del estudio.

Tabla 6-3. Características de los estudios experimentales

Característica	Descripción
Asignación	Controlada por el investigador
Comparación	Grupo intervención y control

Nota. Elementos distintivos de los estudios experimentales en investigación sanitaria.

Los ensayos clínicos permiten evaluar relaciones causales con mayor grado de certeza que los estudios observacionales. No obstante, su implementación requiere consideraciones éticas, logísticas y económicas significativas, así como un análisis estadístico cuidadoso que contemple pérdidas de seguimiento y adherencia al tratamiento.

6.1.4 Comparación entre estudios descriptivos, analíticos y experimentales

Cada tipo de estudio epidemiológico cumple una función específica dentro del proceso de investigación sanitaria. Los estudios descriptivos permiten identificar y caracterizar problemas de salud, los estudios analíticos evalúan asociaciones y generan evidencia explicativa, y los estudios experimentales prueban intervenciones bajo condiciones controladas (Rothman et al., 2008).

Desde una perspectiva bioestadística, la elección del diseño adecuado depende de la pregunta de investigación, la factibilidad del estudio y las consideraciones éticas, reconociendo que no existe un diseño universalmente superior, sino uno más adecuado para cada objetivo científico.



Figura 6-1. Clasificación de los estudios epidemiológicos

6.1.5 Importancia del diseño epidemiológico en la investigación sanitaria

El diseño epidemiológico constituye la base sobre la cual se construye el análisis estadístico y la interpretación de los resultados. Un diseño inapropiado no puede ser corregido mediante técnicas estadísticas avanzadas, lo que subraya la importancia de una planificación metodológica rigurosa desde las etapas iniciales de la investigación (Gordis, 2014).

La bioestadística aplicada, integrada al diseño epidemiológico, permite optimizar la recolección de datos, seleccionar análisis apropiados y garantizar conclusiones válidas y útiles para la toma de decisiones en salud pública y práctica clínica.

6.2 Medidas de frecuencia: incidencia y prevalencia

Las medidas de frecuencia constituyen el núcleo descriptivo de la epidemiología cuantitativa y representan el primer nivel de análisis en la investigación sanitaria. Su función principal es cuantificar la ocurrencia de eventos relacionados con la salud, permitiendo describir la magnitud de las enfermedades, comparar poblaciones y monitorear cambios temporales. Desde la bioestadística aplicada, estas medidas proporcionan una base objetiva para la toma de decisiones en salud pública, la planificación de servicios sanitarios y la evaluación de intervenciones preventivas y terapéuticas (Gordis, 2014; Rothman et al., 2008).

Entre las medidas de frecuencia, la incidencia y la prevalencia ocupan un lugar central debido a su amplia aplicabilidad y a su complementariedad conceptual. Aunque ambas cuantifican la ocurrencia de enfermedad, difieren sustancialmente en su interpretación epidemiológica y en los contextos en los que resultan más apropiadas.

6.2.1 Concepto epidemiológico de incidencia

La incidencia se refiere a la ocurrencia de casos nuevos de una enfermedad o evento de salud en una población definida durante un periodo determinado. Esta medida captura la dinámica de aparición del evento y se interpreta como una estimación del riesgo de enfermar. Desde el punto de vista epidemiológico, la incidencia es esencial para comprender los mecanismos de causalidad y para evaluar la influencia de factores de riesgo y de intervenciones preventivas.

En términos bioestadísticos, la incidencia requiere una definición clara de la población en riesgo, es decir, individuos que están libres de la enfermedad al inicio del periodo de observación y que pueden desarrollar el evento durante el seguimiento. Una definición imprecisa de la población en riesgo conduce a estimaciones sesgadas y a interpretaciones incorrectas de los resultados.

La incidencia resulta particularmente relevante en estudios longitudinales, como los estudios de cohortes, donde se observa la aparición del evento a lo largo del tiempo y se analizan relaciones temporales entre exposición y desenlace.

6.2.2 Incidencia acumulada como medida de riesgo

La incidencia acumulada expresa la proporción de individuos que desarrollan la enfermedad durante un periodo específico, asumiendo que todos los sujetos son seguidos durante el mismo intervalo de tiempo. Se interpreta directamente como un riesgo y suele expresarse como un porcentaje o una proporción (Gordis, 2014).

Desde la bioestadística aplicada, esta medida es adecuada cuando:

- La población es cerrada o relativamente estable.
- El periodo de seguimiento es corto o uniforme.
- Las pérdidas de seguimiento son mínimas.

Su formulación matemática es:

$$\text{Incidencia acumulada} = \frac{\text{Número de casos nuevos}}{\text{Población en riesgo al inicio}}$$

La simplicidad de esta medida facilita su interpretación clínica, aunque su uso debe limitarse a contextos donde los supuestos mencionados se cumplan razonablemente.

Tabla 6-4. Elementos conceptuales de la incidencia acumulada

Elemento	Descripción epidemiológica
Casos nuevos	Aparición del evento
Periodo definido	Marco temporal
Riesgo	Probabilidad de enfermar

Nota. Componentes que permiten interpretar la incidencia acumulada como una medida directa de riesgo poblacional.

6.2.3 Tasa de incidencia y tiempo-persona

La tasa de incidencia incorpora explícitamente el tiempo de observación, lo que permite manejar poblaciones dinámicas y seguimientos desiguales. Esta medida se basa en el concepto de tiempo-persona, que representa la suma del tiempo durante el cual cada individuo permanece en riesgo de desarrollar el evento.

La expresión estadística es:

$$\text{Tasa de incidencia} = \frac{\text{Casos nuevos}}{\text{Tiempo-persona total}}$$

Desde una perspectiva bioestadística, esta formulación ofrece mayor precisión cuando existen pérdidas de seguimiento, entradas tardías o distintos tiempos de observación, situaciones frecuentes en estudios epidemiológicos reales.

La tasa de incidencia permite comparar poblaciones con diferentes duraciones de seguimiento y constituye una herramienta en la evaluación de programas de prevención y en el análisis de enfermedades infecciosas.

6.2.4 Concepto epidemiológico de prevalencia

La prevalencia cuantifica la proporción de individuos que presentan una enfermedad o condición de salud en un momento o periodo determinado, independientemente del momento de inicio del evento. Esta medida refleja la carga total de enfermedad en la población y resulta especialmente útil para la

planificación de recursos sanitarios y la evaluación de necesidades asistenciales.

Desde la bioestadística aplicada, la prevalencia se interpreta como una fotografía del estado de salud poblacional, más que como una medida de riesgo. Su expresión general es:

$$\text{Prevalencia} = \frac{\text{Casos existentes}}{\text{Población total}}$$

La prevalencia se utiliza ampliamente en estudios transversales y encuestas de salud, donde se evalúa la magnitud de enfermedades crónicas y factores de riesgo.

6.2.5 Prevalencia puntual y prevalencia de periodo

La prevalencia puntual se refiere a la proporción de casos existentes en un momento específico, mientras que la prevalencia de periodo incluye todos los casos que se presentan durante un intervalo definido. Esta distinción es relevante desde el punto de vista epidemiológico y estadístico, ya que la prevalencia de periodo suele ser mayor al incorporar casos incidentes y persistentes.

Desde la bioestadística aplicada, la selección del tipo de prevalencia depende del objetivo del estudio. La prevalencia puntual es adecuada para evaluar la carga inmediata de enfermedad, mientras que la prevalencia de periodo ofrece una visión más amplia del impacto sanitario en un intervalo temporal.

Tabla 6-5. Comparación entre tipos de prevalencia

Tipo de prevalencia	Enfoque temporal	Uso principal
Puntual	Momento específico	Diagnóstico situacional
De periodo	Intervalo definido	Evaluación de carga

Nota. Diferencias conceptuales entre prevalencia puntual y prevalencia de periodo en estudios epidemiológicos.

6.2.6 Relación conceptual entre incidencia y prevalencia

La relación entre incidencia y prevalencia constituye uno de los principios fundamentales de la epidemiología. De forma aproximada, puede expresarse como:

$$\text{Prevalencia} = \text{Incidencia} \times \text{Duración promedio}$$

Esta relación permite comprender por qué enfermedades crónicas, aun con baja incidencia, pueden generar alta prevalencia, mientras que enfermedades agudas de corta duración pueden presentar baja prevalencia pese a una incidencia elevada.

Desde la bioestadística aplicada, esta relación subraya la importancia de interpretar ambas medidas de manera conjunta, evitando conclusiones parciales sobre la magnitud real de un problema de salud.

6.2.7 Aplicaciones sanitarias y limitaciones

La incidencia es particularmente útil para:

- Estudiar causas y factores de riesgo.
- Evaluar intervenciones preventivas.
- Analizar la dinámica de transmisión de enfermedades.

La prevalencia, en cambio, resulta más apropiada para:

- Planificar servicios de salud.
- Estimar demanda asistencial.
- Evaluar impacto de enfermedades crónicas.

No obstante, ambas medidas presentan limitaciones. La incidencia requiere seguimiento y puede verse afectada por pérdidas de observación, mientras que la prevalencia puede subestimar enfermedades de corta duración o sobreestimar aquellas con mayor supervivencia.

6.3 Medidas de asociación: riesgo relativo y odds ratio

Las medidas de asociación permiten cuantificar la relación existente entre una exposición y un evento de salud, constituyendo el núcleo explicativo del análisis epidemiológico. Mientras que las medidas de frecuencia describen cuántos casos ocurren en una población, las medidas de asociación responden a una pregunta más compleja: en qué medida la exposición a un factor incrementa o reduce la probabilidad de presentar un desenlace sanitario. Desde la bioestadística aplicada, estas medidas proporcionan una estimación cuantitativa de la fuerza de la relación entre exposición y enfermedad, facilitando la interpretación causal dentro de los límites del diseño del estudio.

Entre las medidas de asociación más utilizadas en epidemiología se encuentran el riesgo relativo y el odds ratio, cada una con aplicaciones específicas según el tipo de estudio y la naturaleza de los datos.

6.3.1 Concepto epidemiológico de asociación

Una asociación epidemiológica existe cuando la frecuencia de un evento de salud difiere entre grupos con distinta exposición a un factor determinado. Desde el punto de vista estadístico, esta diferencia se expresa mediante razones o cocientes que comparan la ocurrencia del evento entre grupos expuestos y no expuestos. La magnitud de la asociación permite evaluar si la exposición se comporta como un posible factor de riesgo o como un factor protector.

La interpretación de estas medidas debe realizarse siempre considerando el diseño del estudio, la presencia de sesgos y factores de confusión, y la plausibilidad biológica de la relación observada.

6.3.2 Riesgo relativo

El riesgo relativo es la medida de asociación más intuitiva y se utiliza principalmente en estudios de cohortes y ensayos clínicos, donde es posible estimar directamente la incidencia del evento en grupos expuestos y no expuestos. El riesgo relativo compara la probabilidad de enfermar en el grupo expuesto con la probabilidad correspondiente en el grupo no expuesto.

Su expresión matemática es:

$$RR = \frac{\text{Riesgo en expuestos}}{\text{Riesgo en no expuestos}}$$

Desde una perspectiva bioestadística, el riesgo relativo permite interpretar de manera directa cuánto se incrementa o reduce el riesgo asociado a una exposición determinada.

- Un valor de RR igual a 1 indica ausencia de asociación.
- Un valor mayor que 1 sugiere un aumento del riesgo asociado a la exposición.
- Un valor menor que 1 indica un efecto protector.

Tabla 6-6. Interpretación general del riesgo relativo

Valor de RR	Interpretación
RR = 1	Sin asociación
RR > 1	Factor de riesgo
RR < 1	Factor protector

Nota. Interpretación básica del riesgo relativo en estudios epidemiológicos analíticos.

El riesgo relativo es ampliamente utilizado en epidemiología clínica debido a su claridad interpretativa. Sin embargo, su estimación requiere seguimiento temporal y una definición precisa de la población en riesgo, lo que limita su uso a ciertos diseños de estudio.

6.3.3 Odds ratio

El odds ratio es una medida de asociación basada en la razón de probabilidades y se emplea principalmente en estudios de casos y controles, donde no es posible calcular directamente la incidencia. Desde el punto de vista bioestadístico, el odds ratio compara las probabilidades de exposición entre casos y controles, proporcionando una estimación indirecta de la asociación entre exposición y enfermedad.

La expresión matemática del odds ratio es:

$$OR = \frac{\text{odds de exposición en casos}}{\text{odds de exposición en controles}}$$

En términos de una tabla de contingencia 2×2, el odds ratio se calcula como:

$$OR = \frac{ad}{bc}$$

donde *a*, *b*, *c* y *d* representan las frecuencias observadas en cada celda de la tabla.

Tabla 6-7. Tabla 2×2 para el cálculo del odds ratio

	Enfermedad	No enfermedad
Expuestos	a	b
No expuestos	c	d

Nota. Estructura básica para el cálculo del odds ratio en estudios de casos y controles.

6.3.4 Interpretación epidemiológica del odds ratio

La interpretación del odds ratio sigue principios similares al riesgo relativo:

- OR igual a 1 indica ausencia de asociación.
- OR mayor que 1 indica asociación positiva entre exposición y enfermedad.
- OR menor que 1 sugiere efecto protector.

Desde el punto de vista epidemiológico, el odds ratio se aproxima al riesgo relativo cuando la enfermedad es poco frecuente en la población. Esta propiedad explica su amplio uso como estimador del riesgo relativo en estudios de casos y controles (Gordis, 2014).

6.3.5 Comparación entre riesgo relativo y odds ratio

Aunque ambas medidas cuantifican asociaciones, difieren en su interpretación y en los diseños de estudio donde se aplican. El riesgo relativo se interpreta directamente como una razón de riesgos, mientras que el odds ratio expresa

una razón de probabilidades, lo que puede resultar menos intuitivo para la práctica clínica.

Tabla 6-8. Comparación entre riesgo relativo y odds ratio

Característica	Riesgo relativo	Odds ratio
Tipo de estudio	Cohortes, ensayos	Casos y controles
Interpretación	Directa	Indirecta
Incidencia requerida	Sí	No

Nota. Diferencias metodológicas entre riesgo relativo y odds ratio según el diseño del estudio.

6.3.6 Intervalos de confianza en medidas de asociación

Desde la bioestadística aplicada, la interpretación del riesgo relativo y del odds ratio debe complementarse con intervalos de confianza, que reflejan la precisión de la estimación. Un intervalo que incluye el valor 1 indica ausencia de asociación estadísticamente significativa, mientras que un intervalo completamente por encima o por debajo de 1 sugiere una asociación consistente (Altman, 1991).

El uso de intervalos de confianza permite evaluar la incertidumbre inherente a las estimaciones y evita interpretaciones basadas únicamente en valores puntuales.

6.3.7 Importancia en salud pública y práctica clínica

Las medidas de asociación son herramientas esenciales para identificar factores de riesgo, priorizar intervenciones y fundamentar políticas de salud pública. Desde la epidemiología clínica, el riesgo relativo orienta decisiones preventivas y terapéuticas, mientras que el odds ratio permite investigar asociaciones en contextos donde el seguimiento no es factible.

La correcta interpretación de estas medidas exige integrar el análisis estadístico con el diseño del estudio, la plausibilidad biológica y el contexto poblacional, principios que constituyen la base de la epidemiología moderna.

6.4 Sesgos y factores de confusión en la investigación sanitaria

En investigación epidemiológica, la validez de los resultados depende no solo del tamaño muestral o de la sofisticación del análisis estadístico, sino de la calidad metodológica con que se recolectan y comparan los datos. Los sesgos y los factores de confusión constituyen amenazas centrales para la validez interna, ya que pueden distorsionar la estimación de la asociación entre exposición y desenlace. Desde la bioestadística aplicada, su identificación, prevención y control forman parte del análisis crítico indispensable para interpretar hallazgos de estudios observacionales y experimentales.

Un resultado estadísticamente significativo puede ser incorrecto si se origina por sesgo sistemático o por confusión no controlada. Por ello, el análisis epidemiológico requiere integrar diseño, recolección de datos y modelado estadístico con criterios rigurosos.

6.4.1 Concepto de sesgo

El sesgo se define como un error sistemático que produce una estimación incorrecta del efecto o asociación de interés. A diferencia del error aleatorio, el sesgo no se reduce aumentando el tamaño muestral, ya que proviene de fallas en el diseño, en la selección de participantes o en la medición de variables.

Desde la perspectiva bioestadística, el sesgo altera el valor esperado de la estimación y puede producir asociaciones falsas, exagerar asociaciones reales o incluso invertir su dirección. Esta característica explica por qué la prevención del sesgo se prioriza en la fase de diseño, antes del análisis.

6.4.2 Sesgo de selección

El sesgo de selección ocurre cuando la manera de seleccionar participantes o la forma de retenerlos en el estudio depende simultáneamente de la exposición y del desenlace, generando una muestra que no representa adecuadamente a la población objetivo. Este sesgo es frecuente en estudios de casos y controles, cuando la selección de controles no refleja la distribución real de exposición en la población fuente.

En estudios de cohortes, el sesgo de selección puede presentarse por pérdidas diferenciales de seguimiento, donde la probabilidad de abandonar el estudio se

asocia con la exposición y con el riesgo de presentar el evento. Desde el punto de vista estadístico, estas pérdidas pueden conducir a estimaciones sesgadas de incidencia y riesgo relativo.

Tabla 6-9. Sesgo de selección en estudios epidemiológicos

Situación	Cómo se origina	Efecto probable
Controles no comparables	Selección inadecuada	OR distorsionado
Pérdidas diferenciales	Abandono relacionado	RR sesgado

Nota. Formas frecuentes de sesgo de selección y su impacto en estimaciones de asociación en estudios analíticos.

6.4.3 Sesgo de información

El sesgo de información, también llamado sesgo de medición o de clasificación, se produce cuando la información sobre la exposición o el desenlace se mide con error sistemático. Este sesgo puede afectar tanto a variables clínicas como a variables reportadas por el paciente, historiales médicos o registros administrativos.

En estudios de casos y controles, un ejemplo típico es el sesgo de recuerdo, donde los casos tienden a reportar exposiciones pasadas de manera distinta a los controles. En cohortes, puede ocurrir sesgo de observación cuando el seguimiento clínico es más intenso en el grupo expuesto, incrementando la probabilidad de detectar el desenlace.

Desde la bioestadística aplicada, el sesgo de información puede clasificarse como diferencial o no diferencial. La clasificación no diferencial tiende a diluir la asociación, mientras que la clasificación diferencial puede producir sesgos impredecibles en magnitud y dirección.

Tabla 6-10. Tipos de sesgo de información

Tipo	Definición	Consecuencia frecuente
No diferencial	Error similar en grupos	Atenúa asociación
Diferencial	Error distinto en grupos	Distorsión variable

Nota. Tipos de sesgo de información según el patrón de error de medición en grupos comparados.

6.4.4 Confusión: concepto y criterios

La confusión ocurre cuando la asociación entre una exposición y un desenlace se mezcla con el efecto de una tercera variable que está relacionada con ambos. Esta tercera variable se denomina confusora y debe cumplir criterios epidemiológicos específicos: asociarse con la exposición, asociarse con el desenlace y no encontrarse en la cadena causal entre exposición y evento.

En investigación sanitaria, variables como edad, sexo, comorbilidades, nivel socioeconómico y estilos de vida suelen actuar como confusoras. La confusión puede crear una asociación aparente sin relación causal o modificar la magnitud real del efecto.

6.4.5 Estrategias de control de confusión en el diseño

La confusión debe controlarse preferentemente en la etapa de diseño. Entre las estrategias principales se encuentran:

- **Aleatorización**, utilizada en ensayos clínicos para distribuir confusores conocidos y desconocidos entre los grupos.
- **Restricción**, que limita la participación a individuos con determinadas características, reduciendo variabilidad, aunque disminuye la generalización.
- **Emparejamiento**, frecuente en casos y controles, donde cada caso se empareja con controles similares respecto a variables como edad o sexo.

Estas estrategias mejoran la comparabilidad de grupos antes del análisis estadístico.

6.4.6 Estrategias de control de confusión en el análisis

Cuando la confusión persiste o no puede controlarse completamente en el diseño, se aplican métodos estadísticos en la etapa analítica. Las estrategias más utilizadas incluyen:

- **Estratificación**, donde se calcula la medida de asociación en subgrupos definidos por el confusor.

- **Modelos multivariantes**, como regresión logística o regresión de Cox, que permiten estimar medidas de asociación ajustadas.

Desde la bioestadística aplicada, estos métodos permiten separar el efecto de la exposición del efecto del confusor, mejorando la validez de la inferencia.

Tabla 6-11. Control de confusión: diseño y análisis

Etapas	Estrategias	Comentario metodológico
Diseño	Aleatorización, restricción, emparejamiento	Prevención temprana
Análisis	Estratificación, modelos multivariantes	Ajuste estadístico

Nota. Estrategias de control de confusión según la fase metodológica del estudio.

6.4.7 Diferencia entre confusión y modificación de efecto

En epidemiología, es importante distinguir confusión de modificación de efecto. La confusión es una distorsión que debe controlarse, mientras que la modificación de efecto describe una situación en la que la magnitud de la asociación varía según niveles de una tercera variable, lo cual representa un hallazgo epidemiológico relevante.

Desde la bioestadística aplicada, la modificación de efecto se identifica cuando las medidas de asociación difieren sustancialmente entre estratos y esta diferencia tiene interpretación clínica o biológica.

6.5 Pruebas estadísticas aplicadas a estudios de cohortes y casos-control

El análisis estadístico en estudios epidemiológicos analíticos constituye una etapa decisiva para evaluar asociaciones entre exposiciones y desenlaces de salud. En particular, los estudios de cohortes y los estudios de casos y controles requieren la aplicación de pruebas estadísticas específicas que permitan determinar si las diferencias observadas entre grupos son atribuibles al azar o reflejan una asociación consistente. Desde la bioestadística aplicada, estas pruebas se seleccionan en función del diseño del estudio, del tipo de variables analizadas y del tamaño muestral disponible.

La correcta elección e interpretación de las pruebas estadísticas es esencial para garantizar conclusiones válidas y evitar inferencias erróneas en investigación sanitaria.

6.5.1 Fundamento estadístico en estudios analíticos

Los estudios de cohortes y casos-control se basan en la comparación de grupos definidos por su exposición o por su condición de enfermedad. Desde el punto de vista estadístico, esta comparación implica contrastar proporciones, tasas o razones de ocurrencia entre grupos, lo que exige herramientas inferenciales adecuadas.

En estudios de cohortes, donde se estima la incidencia, el análisis se centra en la comparación de riesgos. En estudios de casos y controles, donde la incidencia no es directamente observable, el análisis se enfoca en la comparación de probabilidades de exposición. En ambos escenarios, las pruebas estadísticas permiten evaluar la significación de las asociaciones observadas y cuantificar la incertidumbre asociada a las estimaciones (Altman, 1991).

6.5.2 Prueba Chi-cuadrado en estudios epidemiológicos

La prueba Chi-cuadrado es una de las herramientas más utilizadas para evaluar la asociación entre variables cualitativas en estudios epidemiológicos. En estudios de cohortes y casos y controles, esta prueba permite determinar si existe una diferencia estadísticamente significativa entre las proporciones observadas en los grupos comparados (Pagano y Gauvreau, 2018).

Desde la bioestadística aplicada, la prueba Chi-cuadrado se utiliza cuando:

- Las variables son categóricas.
- Las observaciones son independientes.
- Los valores esperados en las celdas de la tabla son suficientes para garantizar validez estadística.

Su interpretación se basa en la comparación entre el valor calculado del estadístico y su distribución teórica, permitiendo evaluar la evidencia estadística de asociación.

Tabla 6-12. Uso de la prueba Chi-cuadrado en estudios analíticos

Aspecto	Aplicación
Tipo de variables	Cualitativas
Diseño	Cohortes y casos-control
Objetivo	Evaluar asociación

Nota. Aplicación general de la prueba Chi-cuadrado en estudios epidemiológicos analíticos.

6.5.3 Prueba exacta de Fisher

Cuando el tamaño muestral es reducido o existen frecuencias esperadas bajas en las celdas de la tabla de contingencia, la prueba exacta de Fisher constituye una alternativa adecuada a la prueba Chi-cuadrado. Desde el punto de vista bioestadístico, esta prueba calcula la probabilidad exacta de observar la distribución de datos bajo la hipótesis de ausencia de asociación (Altman, 1991).

La prueba exacta de Fisher es especialmente relevante en estudios de casos y controles con muestras pequeñas, donde la aproximación de la Chi-cuadrado puede resultar inapropiada.

6.5.4 Pruebas para comparación de tasas e incidencias

En estudios de cohortes, el análisis puede centrarse en la comparación de tasas de incidencia entre grupos expuestos y no expuestos. Desde la bioestadística aplicada, estas comparaciones se realizan mediante razones de tasas y pruebas basadas en distribuciones de Poisson, especialmente cuando los eventos son relativamente raros (Rothman et al., 2008).

Estas pruebas permiten evaluar diferencias en la ocurrencia de eventos ajustadas por tiempo de observación, fortaleciendo la interpretación epidemiológica de los resultados.

6.5.5 **Regresión logística en estudios de casos y controles**

La regresión logística es una herramienta estadística ampliamente utilizada en estudios de casos y controles para evaluar la asociación entre una o varias exposiciones y un desenlace binario. Desde la bioestadística aplicada, este modelo permite estimar odds ratios ajustados, controlando simultáneamente múltiples variables potencialmente confusoras.

La regresión logística resulta particularmente valiosa cuando se analizan exposiciones múltiples o cuando se requiere ajustar por variables sociodemográficas, clínicas o ambientales.

Tabla 6-13. Aplicaciones de la regresión logística

Característica	Utilidad
Variable dependiente	Binaria
Resultado	Odds ratio ajustado
Control	Confusión

Nota. Principales aplicaciones de la regresión logística en estudios epidemiológicos analíticos.

6.5.6 **Intervalos de confianza y significación estadística**

En estudios epidemiológicos, la significación estadística debe interpretarse conjuntamente con los intervalos de confianza de las medidas de asociación. Desde la bioestadística aplicada, los intervalos de confianza aportan información sobre la precisión de la estimación y la magnitud plausible del efecto observado (Altman, 1991).

Un intervalo de confianza que no incluye el valor nulo indica evidencia estadística de asociación, mientras que la amplitud del intervalo refleja el grado de incertidumbre asociado al tamaño muestral y a la variabilidad de los datos.

6.5.7 **Consideraciones metodológicas y limitaciones**

La aplicación de pruebas estadísticas en estudios de cohortes y casos y controles debe considerar cuidadosamente los supuestos de cada prueba, la

calidad de los datos y el diseño del estudio. Un uso inapropiado de las pruebas estadísticas puede conducir a conclusiones erróneas, aun cuando los cálculos sean correctos desde el punto de vista matemático (Rothman et al., 2014).

Desde la epidemiología aplicada, la interpretación de los resultados estadísticos debe integrarse con el conocimiento clínico, la plausibilidad biológica y el contexto poblacional, evitando inferencias exclusivamente basadas en valores de significación.

6.6 Meta-análisis y síntesis de resultados de investigación

El meta-análisis representa una herramienta avanzada del análisis epidemiológico que permite integrar cuantitativamente los resultados de múltiples estudios independientes sobre una misma pregunta de investigación. Desde la bioestadística aplicada, el meta-análisis no se limita a una revisión narrativa, sino que utiliza métodos estadísticos formales para combinar estimaciones de efecto, aumentar la precisión de los resultados y mejorar la potencia estadística global de la evidencia disponible (Borenstein et al., 2009; Higgins et al., 2019).

En el ámbito de las ciencias de la salud, el meta-análisis es ampliamente utilizado para evaluar la efectividad de intervenciones clínicas, identificar factores de riesgo, comparar tratamientos y fundamentar guías de práctica clínica y políticas sanitarias.

6.6.1 Concepto y fundamentos del meta-análisis

El meta-análisis puede definirse como un conjunto de procedimientos estadísticos destinados a combinar cuantitativamente los resultados de estudios primarios que abordan una misma hipótesis o efecto. Cada estudio aporta una estimación puntual del efecto y una medida de variabilidad, que son ponderadas en función de su precisión para obtener una estimación combinada más robusta.

Desde la perspectiva epidemiológica, el meta-análisis se apoya en el supuesto de que los estudios incluidos son suficientemente comparables en términos de diseño, población, exposición y desenlace. La bioestadística aplicada permite evaluar esta comparabilidad y cuantificar la variabilidad entre estudios mediante indicadores específicos de heterogeneidad.

6.6.2 Tipos de medidas de efecto en meta-análisis

Las medidas de efecto utilizadas en meta-análisis dependen del tipo de desenlace analizado. En estudios epidemiológicos y clínicos, las más frecuentes incluyen el riesgo relativo, el odds ratio, la diferencia de medias y la diferencia de medias estandarizada.

Cada estudio contribuye con su estimación del efecto y su error estándar, lo que permite asignar mayor peso a los estudios con menor variabilidad. Este principio garantiza que la estimación combinada refleje de manera equilibrada la evidencia disponible.

Tabla 6-14. Medidas de efecto utilizadas en meta-análisis

Tipo de desenlace	Medida de efecto	Uso habitual
Binario	Riesgo relativo, odds ratio	Estudios clínicos
Continuo	Diferencia de medias	Variables cuantitativas
Continuo heterogéneo	Diferencia estandarizada	Escalas distintas

Nota. Principales medidas de efecto empleadas en meta-análisis según la naturaleza del desenlace.

6.6.3 Modelos estadísticos en meta-análisis

Desde la bioestadística aplicada, los modelos de meta-análisis se clasifican principalmente en modelos de efecto fijo y modelos de efectos aleatorios. El modelo de efecto fijo asume que todos los estudios estiman un único efecto verdadero, y que las diferencias observadas se deben exclusivamente al error aleatorio. En contraste, el modelo de efectos aleatorios considera que el efecto puede variar entre estudios debido a diferencias metodológicas o poblacionales.

La elección del modelo adecuado depende del grado de heterogeneidad entre los estudios, evaluado mediante estadísticas específicas.

6.6.4 Evaluación de la heterogeneidad

La heterogeneidad se refiere a la variabilidad entre los resultados de los estudios incluidos más allá de lo esperado por azar. Desde el punto de vista epidemiológico y estadístico, la heterogeneidad puede deberse a diferencias en diseño, población, intervención o medición del desenlace (Higgins et al., 2022).

Los indicadores más utilizados para cuantificar la heterogeneidad son el estadístico Q y el índice I², que expresa el porcentaje de variabilidad total atribuible a diferencias reales entre estudios.

$$I^2 = \frac{Q - df}{Q} \times 100$$

Valores elevados de I² indican mayor heterogeneidad y justifican el uso de modelos de efectos aleatorios y análisis exploratorios adicionales.

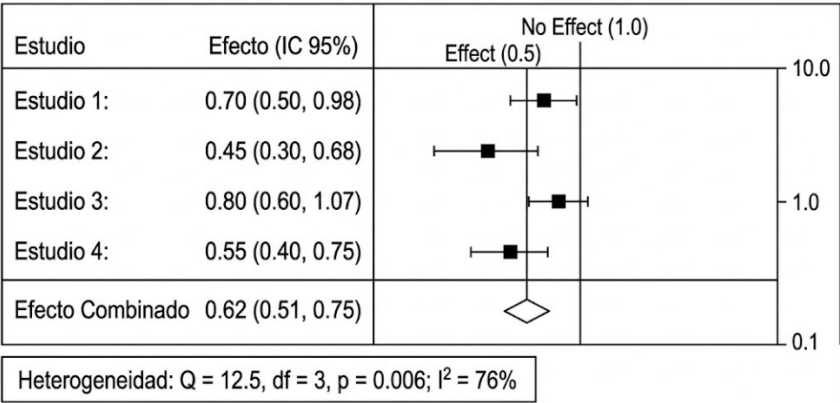


Figura 6-2. Diagrama de bosque en un meta-análisis. Representación gráfica de resultados individuales y efecto combinado en un meta-análisis.

6.6.5 Síntesis cuantitativa y visualización de resultados

La síntesis de resultados en meta-análisis combina la estimación puntual y el intervalo de confianza de cada estudio para obtener una medida global del efecto. Desde la bioestadística aplicada, esta síntesis se representa visualmente

mediante diagramas de bosque, que facilitan la comparación entre estudios y la evaluación de consistencia de los resultados (Borenstein et al., 2009).

Estos diagramas permiten identificar estudios con efectos extremos, evaluar la dirección general de la asociación y comunicar los resultados de forma clara a clínicos y responsables de salud pública.

6.6.6 Sesgo de publicación y validez de la síntesis

Un aspecto crítico del meta-análisis es el sesgo de publicación, que ocurre cuando los estudios con resultados estadísticamente significativos tienen mayor probabilidad de ser publicados. Este sesgo puede distorsionar la estimación del efecto combinado. Desde la bioestadística aplicada, se utilizan métodos gráficos y pruebas estadísticas para evaluar este fenómeno, contribuyendo a una interpretación más prudente de los resultados (Higgins et al., 2022).

La identificación del sesgo de publicación es esencial para preservar la validez externa de la síntesis de evidencia.

6.6.7 Importancia del meta-análisis en la investigación sanitaria

El meta-análisis ocupa un lugar central en la jerarquía de evidencia científica, ya que permite integrar resultados dispersos y ofrecer conclusiones más precisas que los estudios individuales. En epidemiología y medicina clínica, esta herramienta sustenta la elaboración de guías clínicas, evaluaciones de tecnologías sanitarias y decisiones de política pública (Borenstein et al., 2009).

Desde la bioestadística aplicada, el meta-análisis ejemplifica la integración entre rigor metodológico, análisis estadístico avanzado y utilidad práctica de la investigación sanitaria.

6.7 Aplicaciones prácticas en vigilancia epidemiológica y salud pública

La vigilancia epidemiológica constituye una de las aplicaciones más relevantes del análisis epidemiológico y de la bioestadística aplicada en el ámbito de la salud pública. Su finalidad es recolectar, analizar, interpretar y difundir de manera sistemática datos relacionados con eventos de salud, con el propósito de orientar la toma de decisiones, la planificación de intervenciones y la

evaluación de políticas sanitarias. Desde una perspectiva cuantitativa, la vigilancia epidemiológica integra medidas de frecuencia, medidas de asociación, control de sesgos y síntesis de evidencia para anticipar riesgos y responder oportunamente a problemas de salud poblacional (Bonita et al., 2012; Gordis, 2014).

6.7.1 Vigilancia epidemiológica como proceso cuantitativo

La vigilancia epidemiológica moderna se sustenta en métodos estadísticos que permiten transformar datos en información útil para la acción sanitaria. Este proceso incluye la definición de eventos de interés, la recolección sistemática de datos, el análisis estadístico continuo y la retroalimentación a los tomadores de decisiones.

Desde la bioestadística aplicada, la vigilancia se apoya principalmente en:

- Medidas de incidencia y prevalencia para monitorear la magnitud de los eventos.
- Análisis de tendencias temporales para detectar cambios inusuales.
- Comparaciones espaciales para identificar agrupamientos o focos de riesgo.

La calidad del sistema de vigilancia depende directamente de la precisión de los datos y de la adecuada selección de indicadores epidemiológicos.

Tabla 6-15. Componentes cuantitativos de la vigilancia epidemiológica

Componente	Herramienta bioestadística
Magnitud del evento	Incidencia, prevalencia
Tendencia temporal	Series de tiempo
Comparación poblacional	Medidas de asociación

Nota. Principales componentes cuantitativos utilizados en sistemas de vigilancia epidemiológica en salud pública.

6.7.2 Aplicaciones en salud pública

En salud pública, la bioestadística aplicada permite priorizar problemas sanitarios, asignar recursos y evaluar el impacto de intervenciones. Las medidas epidemiológicas orientan decisiones como la implementación de campañas de vacunación, programas de prevención de enfermedades crónicas y estrategias de control de brotes (Bonita et al., 2006).

El análisis epidemiológico también facilita la identificación de poblaciones vulnerables y la evaluación de inequidades en salud, contribuyendo a políticas basadas en evidencia cuantitativa y criterios de equidad sanitaria.

6.7.3 Uso de indicadores epidemiológicos para la toma de decisiones

Los indicadores epidemiológicos sintetizan información compleja en medidas comprensibles para gestores y decisores. Desde la bioestadística aplicada, su correcta selección e interpretación es esencial para evitar conclusiones erróneas y orientar acciones proporcionales al riesgo observado.

Indicadores como tasas específicas por edad, sexo o región permiten focalizar intervenciones y evaluar su efectividad a lo largo del tiempo.

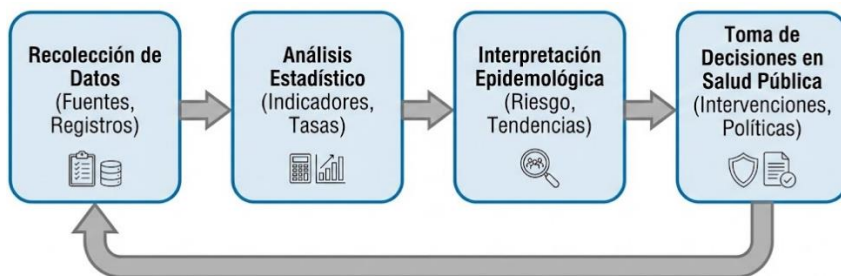


Figura 6-3. Flujo de información en vigilancia epidemiológica. Proceso continuo de análisis y retroalimentación en la vigilancia epidemiológica poblacional.

6.7.4 Integración de estudios epidemiológicos y vigilancia

Los estudios descriptivos, analíticos y experimentales alimentan los sistemas de vigilancia epidemiológica con evidencia estructurada. Los estudios descriptivos identifican patrones iniciales, los analíticos evalúan factores

asociados y los experimentales permiten valorar intervenciones. Desde la bioestadística aplicada, esta integración fortalece la capacidad predictiva y la respuesta sanitaria ante eventos emergentes.

6.7.4.1 Ejercicio

Planteamiento

En una provincia con una población de 120 000 habitantes, el sistema de vigilancia epidemiológica registra durante un año los siguientes datos sobre una enfermedad respiratoria:

- Casos existentes al inicio del año: 1 800
- Casos nuevos durante el año: 2 400
- Población promedio en riesgo: 118 200

Además, se identifica un factor de exposición ambiental presente en una parte de la población. Los datos analíticos se resumen así:

Exposición	Enfermos	No enfermos
Expuestos	1 200	18 800
No expuestos	1 200	97 000

Actividad

1. Calcular la incidencia anual.
2. Calcular la prevalencia al final del año.
3. Estimar el riesgo relativo asociado a la exposición.
4. Interpretar los resultados desde una perspectiva de salud pública.
5. Explicar cómo estos indicadores orientan acciones de vigilancia.

Solución

Incidencia anual

$$\text{Incidencia} = \frac{2400}{118200} = 0.0203$$

Incidencia anual aproximada del 2.0 por ciento.

Prevalencia al final del año

Casos totales al final del año:

$$1800 + 2400 = 4200$$

$$\text{Prevalencia} = \frac{4200}{120000} = 0.035$$

Prevalencia del 3.5 por ciento.

Riesgo relativo

Riesgo en expuestos:

$$\frac{1200}{20000} = 0.06$$

Riesgo en no expuestos:

$$\frac{1200}{98200} = 0.0122$$

$$RR = \frac{0.06}{0.0122} = 4.9$$

Interpretación epidemiológica y sanitaria

La incidencia evidencia una transmisión activa de la enfermedad, mientras que la prevalencia indica una carga sanitaria considerable. El riesgo relativo cercano a 5 sugiere una fuerte asociación entre la exposición ambiental y la enfermedad, lo que justifica intervenciones prioritarias en la población expuesta.

Desde la vigilancia epidemiológica, estos resultados orientan:

- Intensificación de medidas preventivas en zonas expuestas.
- Monitoreo continuo de la incidencia.
- Evaluación del impacto de intervenciones ambientales y sanitarias.

Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud

CAPÍTULO VII

Inteligencia Artificial y Aprendizaje Automático en Bioestadística

AUTOR: Regalado Jalca Julio Johnny



7 Inteligencia Artificial y Aprendizaje Automático en Bioestadística

7.1 Introducción al aprendizaje automático en ciencias de la salud

El aprendizaje automático ha emergido como una de las herramientas más influyentes en la transformación de las ciencias de la salud durante las últimas décadas. Su incorporación en el análisis de datos biomédicos responde a la creciente complejidad, volumen y heterogeneidad de la información generada en contextos clínicos, epidemiológicos y de laboratorio. Desde la perspectiva de la bioestadística aplicada, el aprendizaje automático no sustituye los fundamentos estadísticos clásicos, sino que los complementa y amplía, permitiendo identificar patrones, realizar predicciones y apoyar la toma de decisiones en escenarios donde los métodos tradicionales presentan limitaciones (Hastie et al., 2009; James et al., 2021).

En el ámbito sanitario, los datos provienen de múltiples fuentes, como historias clínicas electrónicas, pruebas de laboratorio, imágenes médicas, registros epidemiológicos y dispositivos de monitoreo continuo. El aprendizaje automático ofrece un conjunto de métodos capaces de aprender de estos datos y generar modelos predictivos o descriptivos con un alto grado de adaptabilidad, lo que resulta especialmente relevante en contextos clínicos caracterizados por la incertidumbre y la variabilidad interindividual.

7.1.1 Relación entre bioestadística y aprendizaje automático

La bioestadística ha sido históricamente la disciplina encargada de analizar datos en ciencias de la salud, centrándose en la inferencia, la estimación y la evaluación de la incertidumbre. El aprendizaje automático, por su parte, se orienta principalmente a la predicción y a la detección de patrones complejos a partir de grandes volúmenes de datos. Ambos enfoques comparten una base matemática común, sustentada en la probabilidad, el álgebra lineal y la optimización, aunque difieren en sus objetivos y supuestos metodológicos.

Desde una perspectiva aplicada, la bioestadística aporta el marco conceptual para definir correctamente las variables, controlar sesgos, interpretar resultados y evaluar la validez clínica de los modelos. El aprendizaje automático, en cambio, permite manejar relaciones no lineales, interacciones

complejas y estructuras de datos de alta dimensión, cada vez más frecuentes en la investigación sanitaria moderna.

7.1.2 Concepto de aprendizaje automático en salud

El aprendizaje automático puede definirse como un conjunto de métodos computacionales que permiten a un sistema mejorar su desempeño en una tarea específica a partir de la experiencia, entendida como datos. En ciencias de la salud, esta experiencia se materializa en registros clínicos, resultados diagnósticos, antecedentes epidemiológicos y desenlaces clínicos observados (Bishop, 2006).

A diferencia de los modelos estadísticos clásicos, en los que la forma funcional suele definirse explícitamente, los algoritmos de aprendizaje automático aprenden la estructura del modelo directamente a partir de los datos. Esta característica resulta particularmente útil en escenarios clínicos donde los mecanismos subyacentes no son completamente conocidos o son difíciles de modelar mediante enfoques paramétricos tradicionales.

7.1.3 Evolución del aprendizaje automático en el ámbito sanitario

La aplicación del aprendizaje automático en salud ha evolucionado de manera progresiva, paralela al desarrollo de la capacidad computacional y a la digitalización de los sistemas sanitarios. En sus primeras etapas, estos métodos se emplearon principalmente en sistemas expertos y reglas basadas en conocimiento clínico. Posteriormente, el acceso a grandes bases de datos permitió el desarrollo de modelos estadísticos avanzados y algoritmos de clasificación y predicción más robustos.

En la actualidad, el aprendizaje automático se utiliza en áreas tan diversas como el diagnóstico asistido por computador, la predicción de riesgo clínico, la estratificación de pacientes, la vigilancia epidemiológica y la medicina personalizada. Este desarrollo ha reforzado la necesidad de integrar conocimientos de bioestadística, informática médica y ciencias de datos dentro de la formación de profesionales de la salud.

7.1.4 Tipos de problemas abordados mediante aprendizaje automático en salud

En ciencias de la salud, el aprendizaje automático se aplica principalmente a problemas de clasificación, regresión, agrupamiento y detección de anomalías. Estos problemas se relacionan directamente con tareas clínicas habituales, como la identificación de pacientes con alto riesgo, la predicción de resultados terapéuticos y la detección temprana de eventos adversos (James et al., 2021).

Desde la bioestadística aplicada, resulta fundamental comprender la naturaleza del problema antes de seleccionar un algoritmo, ya que la interpretación clínica y la validez del modelo dependen de una correcta formulación del objetivo analítico.

Tabla 7-1. Principales problemas abordados por el aprendizaje automático en salud

Tipo de problema	Ejemplo en salud	Objetivo
Clasificación	Diagnóstico positivo o negativo	Decisión clínica
Regresión	Predicción de glucosa	Estimación cuantitativa
Agrupamiento	Fenotipos clínicos	Estratificación
Detección de anomalías	Eventos adversos	Alerta temprana

Nota. Problemas analíticos frecuentes en ciencias de la salud abordados mediante aprendizaje automático.

7.1.5 Datos sanitarios y aprendizaje automático

El desempeño de los modelos de aprendizaje automático depende en gran medida de la calidad de los datos utilizados. En el ámbito sanitario, los datos suelen presentar características particulares, como valores faltantes, errores de medición, desbalance entre clases y alta correlación entre variables. Desde la bioestadística aplicada, estas características deben ser evaluadas y tratadas adecuadamente antes del modelado.

La preparación de los datos, que incluye limpieza, codificación de variables, normalización y selección de características, constituye una etapa crítica del

proceso de aprendizaje automático en salud. Una preparación inadecuada puede conducir a modelos aparentemente precisos, pero clínicamente irrelevantes o sesgados.

7.1.6 Interpretabilidad y toma de decisiones clínicas

Un aspecto central en la aplicación del aprendizaje automático en ciencias de la salud es la interpretabilidad de los modelos. A diferencia de otros dominios, en el contexto sanitario no basta con obtener predicciones precisas; es necesario comprender cómo y por qué el modelo llega a una determinada conclusión, especialmente cuando estas influyen en decisiones clínicas.

Desde la bioestadística aplicada, la interpretabilidad se relaciona con la capacidad de explicar el efecto de las variables, evaluar la coherencia clínica de los resultados y estimar la incertidumbre asociada a las predicciones. Este enfoque refuerza la necesidad de integrar modelos de aprendizaje automático con principios estadísticos sólidos y conocimiento clínico experto.

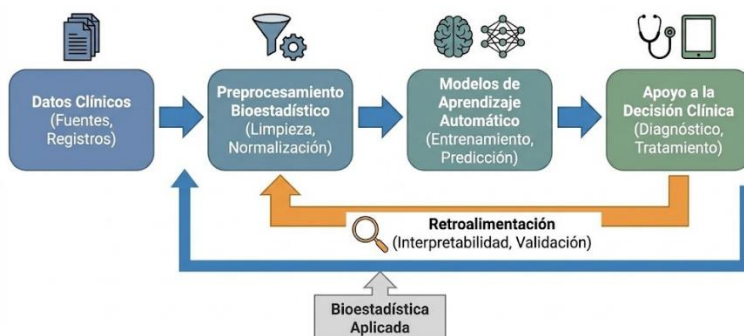


Figura 7-1. Integración del aprendizaje automático en el análisis bioestadístico en salud. Relación entre bioestadística, aprendizaje automático y toma de decisiones en ciencias de la salud.

7.1.7 Relevancia del aprendizaje automático en la formación en ciencias de la salud

La incorporación del aprendizaje automático en la formación de profesionales de la salud responde a la necesidad de comprender y evaluar críticamente los modelos que ya forman parte de la práctica clínica moderna. Desde la

educación en bioestadística, este enfoque permite ampliar la capacidad analítica del profesional, fortaleciendo su rol en equipos interdisciplinarios y su participación en investigaciones basadas en datos.

Comprender los fundamentos del aprendizaje automático, sus alcances y limitaciones, resulta indispensable para garantizar un uso responsable, ético y clínicamente pertinente de estas herramientas en la atención sanitaria.

7.2 Tipos de algoritmos: supervisados, no supervisados y de refuerzo

El aprendizaje automático comprende un conjunto amplio de algoritmos que pueden clasificarse según la forma en que utilizan la información disponible durante el proceso de aprendizaje. En ciencias de la salud, esta clasificación resulta especialmente relevante, ya que los datos clínicos y epidemiológicos presentan estructuras variadas, niveles distintos de rotulación y objetivos analíticos diversos. Desde la bioestadística aplicada, comprender los tipos de algoritmos permite seleccionar métodos adecuados, interpretar correctamente los resultados y evaluar su pertinencia clínica.

De manera general, los algoritmos de aprendizaje automático se agrupan en tres grandes categorías: algoritmos supervisados, algoritmos no supervisados y algoritmos de aprendizaje por refuerzo. Cada uno responde a problemas específicos y presenta ventajas y limitaciones que deben considerarse en el ámbito sanitario.

7.2.1 Aprendizaje supervisado

El aprendizaje supervisado se basa en el uso de conjuntos de datos en los que cada observación incluye tanto las variables predictoras como una variable de salida conocida, denominada etiqueta o respuesta. En ciencias de la salud, estas etiquetas suelen corresponder a diagnósticos clínicos, desenlaces terapéuticos o estados de salud definidos previamente.

Desde la bioestadística aplicada, el aprendizaje supervisado se relaciona estrechamente con los modelos de regresión y clasificación tradicionales, aunque incorpora mayor flexibilidad para capturar relaciones no lineales y efectos complejos entre variables. Este enfoque es ampliamente utilizado en

diagnóstico asistido, predicción de riesgo clínico y pronóstico de enfermedades.

Los problemas abordados mediante aprendizaje supervisado se dividen principalmente en clasificación y regresión. En la clasificación, la variable de salida es categórica, mientras que en la regresión es continua.

Tabla 7-2. Aplicaciones del aprendizaje supervisado en salud

Tipo de problema	Variable de salida	Ejemplo clínico
Clasificación	Categórica	Diagnóstico positivo o negativo
Regresión	Continua	Predicción de presión arterial

Nota. Usos frecuentes del aprendizaje supervisado en contextos clínicos y epidemiológicos.

Entre los algoritmos supervisados más utilizados en salud se encuentran la regresión logística, los árboles de decisión, las máquinas de soporte vectorial, los métodos de ensamble y las redes neuronales. La selección del algoritmo depende del tipo de datos, del tamaño muestral y del equilibrio entre precisión e interpretabilidad.

7.2.2 Aprendizaje no supervisado

El aprendizaje no supervisado se aplica cuando los datos no cuentan con etiquetas conocidas, y el objetivo principal es identificar estructuras, patrones o agrupamientos inherentes a la información. En ciencias de la salud, este enfoque se utiliza para explorar bases de datos complejas, identificar subgrupos de pacientes y generar hipótesis clínicas o epidemiológicas.

Desde la bioestadística aplicada, el aprendizaje no supervisado se relaciona con técnicas exploratorias clásicas, como el análisis de conglomerados y la reducción de dimensionalidad. Sin embargo, los métodos modernos permiten manejar grandes volúmenes de datos y detectar patrones que no son evidentes mediante enfoques tradicionales.

Los algoritmos más representativos del aprendizaje no supervisado incluyen el agrupamiento por k medias, el agrupamiento jerárquico y los métodos de reducción de dimensionalidad como el análisis de componentes principales.

Tabla 7-3. Algoritmos no supervisados y su uso en salud

Algoritmo	Objetivo	Aplicación sanitaria
k medias	Agrupamiento	Fenotipos clínicos
Agrupamiento jerárquico	Estructura de datos	Subtipos de enfermedad
Componentes principales	Reducción dimensional	Datos ómicos

Nota. Principales algoritmos no supervisados utilizados en análisis exploratorio de datos sanitarios.

El aprendizaje no supervisado desempeña un rol fundamental en la medicina personalizada, al permitir la identificación de perfiles clínicos o biológicos que pueden responder de manera distinta a tratamientos específicos.

7.2.3 Aprendizaje por refuerzo

El aprendizaje por refuerzo se diferencia de los enfoques supervisados y no supervisados en que el algoritmo aprende a partir de la interacción con un entorno, recibiendo retroalimentación en forma de recompensas o penalizaciones. En ciencias de la salud, este enfoque se ha explorado principalmente en la optimización de decisiones secuenciales, como la selección adaptativa de tratamientos o la gestión de recursos sanitarios.

Desde la bioestadística aplicada, el aprendizaje por refuerzo introduce un marco dinámico que incorpora la dimensión temporal y la toma de decisiones bajo incertidumbre. Aunque su uso clínico aún es limitado en comparación con otros enfoques, su potencial es considerable en áreas como la dosificación personalizada y la planificación de terapias.

Los elementos fundamentales del aprendizaje por refuerzo incluyen el agente, el entorno, las acciones y la función de recompensa, que guían el proceso de aprendizaje hacia políticas de decisión óptimas.

Tabla 7-4. Componentes del aprendizaje por refuerzo

Componente	Descripción
Agente	Sistema que aprende
Entorno	Contexto clínico
Recompensa	Criterio de desempeño

Nota. Elementos básicos que estructuran el aprendizaje por refuerzo en aplicaciones sanitarias.

7.2.4 Comparación entre los tipos de aprendizaje automático

Cada tipo de aprendizaje automático responde a necesidades analíticas distintas y presenta implicaciones específicas para la interpretación clínica. El aprendizaje supervisado destaca por su utilidad predictiva, el no supervisado por su capacidad exploratoria y el aprendizaje por refuerzo por su enfoque en la toma de decisiones adaptativa (James et al., 2021).

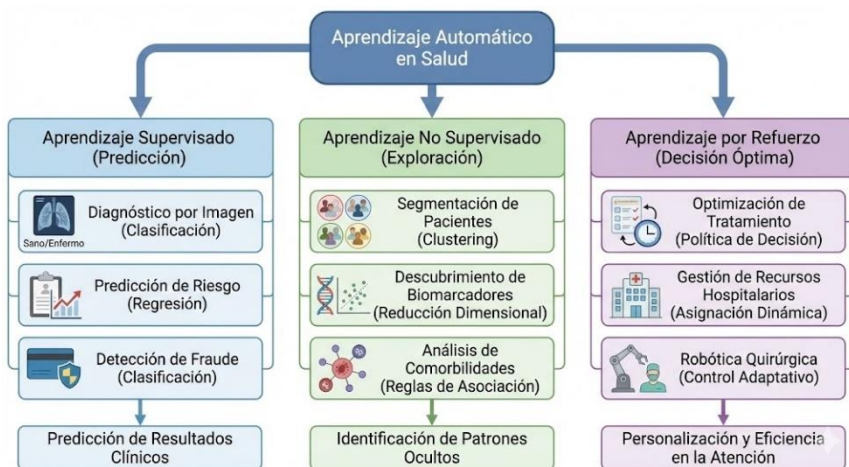


Figura 7-2. Clasificación de algoritmos de aprendizaje automático en salud

Desde la bioestadística aplicada, la elección del enfoque adecuado debe basarse en la disponibilidad de datos etiquetados, el objetivo del análisis y la necesidad de interpretabilidad y validación clínica.

Tabla 7-5. Comparación de tipos de aprendizaje automático

Tipo	Datos etiquetados	Objetivo principal
Supervisado	Sí	Predicción
No supervisado	No	Exploración
Refuerzo	Parcial	Decisión óptima

Nota. Comparación general de enfoques de aprendizaje automático según datos y objetivos analíticos.

7.3 Uso de redes neuronales en diagnóstico médico automatizado

Las redes neuronales artificiales representan uno de los desarrollos más relevantes del aprendizaje automático aplicado a las ciencias de la salud, particularmente en el ámbito del diagnóstico médico automatizado. Su capacidad para modelar relaciones complejas y no lineales entre múltiples variables las convierte en herramientas especialmente adecuadas para el análisis de datos clínicos de alta dimensionalidad, como imágenes médicas, señales biomédicas y grandes volúmenes de información clínica estructurada y no estructurada. Desde la bioestadística aplicada, las redes neuronales deben entenderse no solo como modelos predictivos potentes, sino como sistemas que requieren una validación rigurosa, interpretación clínica y evaluación estadística cuidadosa.

En el contexto sanitario, el diagnóstico automatizado mediante redes neuronales busca apoyar, no reemplazar, el juicio clínico, proporcionando estimaciones objetivas y reproducibles que contribuyan a la detección temprana de enfermedades y a la mejora de la precisión diagnóstica.

7.3.1 Fundamentos conceptuales de las redes neuronales

Las redes neuronales artificiales se inspiran en la estructura y funcionamiento del sistema nervioso biológico. Están compuestas por unidades básicas denominadas neuronas artificiales, organizadas en capas y conectadas mediante pesos ajustables. Cada neurona recibe un conjunto de entradas, realiza una combinación lineal ponderada y aplica una función de activación que introduce no linealidad al modelo.

Desde la bioestadística aplicada, este mecanismo puede interpretarse como una generalización de los modelos de regresión, donde las relaciones entre variables no están restringidas a formas lineales ni a un número reducido de interacciones. Esta flexibilidad permite a las redes neuronales capturar patrones complejos presentes en datos biomédicos reales.

7.3.2 Arquitectura básica de una red neuronal

Una red neuronal típica utilizada en diagnóstico médico automatizado consta de tres tipos de capas: capa de entrada, una o varias capas ocultas y capa de salida. La capa de entrada recibe las variables clínicas o características extraídas de los datos, mientras que las capas ocultas transforman progresivamente la información. La capa de salida produce la predicción final, que puede ser una probabilidad diagnóstica o una clasificación categórica.

El número de capas y neuronas, así como la elección de funciones de activación, influyen directamente en la capacidad del modelo para generalizar a nuevos datos, aspecto crítico en aplicaciones clínicas.

Tabla 7-6. Componentes de una red neuronal artificial

Componente	Función
Neurona	Unidad de procesamiento
Pesos	Importancia de entradas
Función de activación	Introduce no linealidad
Capas	Estructura del modelo

Nota. Elementos estructurales básicos que conforman una red neuronal artificial.

7.3.3 Entrenamiento y aprendizaje en redes neuronales

El entrenamiento de una red neuronal consiste en ajustar los pesos de las conexiones para minimizar una función de error que cuantifica la discrepancia entre las predicciones del modelo y los valores reales observados. Este proceso se realiza mediante algoritmos de optimización iterativa, siendo el descenso del gradiente uno de los métodos más utilizados.

Desde la bioestadística aplicada, el entrenamiento debe realizarse utilizando conjuntos de datos independientes para entrenamiento y validación, con el fin de evitar el sobreajuste. El sobreajuste ocurre cuando el modelo aprende patrones específicos del conjunto de entrenamiento que no se reproducen en datos nuevos, lo que compromete su utilidad clínica.

7.3.4 Aplicaciones de redes neuronales en diagnóstico médico

Las redes neuronales han demostrado un desempeño destacado en múltiples áreas del diagnóstico médico automatizado. En imagenología, se utilizan para la detección de lesiones en radiografías, tomografías computarizadas y resonancias magnéticas. En cardiología, se aplican al análisis de electrocardiogramas y señales fisiológicas. En laboratorio clínico, apoyan la interpretación de perfiles bioquímicos complejos y patrones hematológicos.

Desde la bioestadística aplicada, estas aplicaciones requieren una evaluación rigurosa de sensibilidad, especificidad y valor predictivo, asegurando que el modelo aporte beneficios reales en la práctica clínica.

Tabla 7-7. Aplicaciones diagnósticas de redes neuronales en salud

Área clínica	Tipo de datos	Objetivo diagnóstico
Imagenología	Imágenes médicas	Detección de lesiones
Cardiología	Señales biomédicas	Identificación de arritmias
Laboratorio clínico	Datos numéricos	Clasificación diagnóstica

Nota. Principales áreas clínicas donde se aplican redes neuronales para diagnóstico automatizado.

7.3.5 Redes neuronales profundas en salud

Las redes neuronales profundas, caracterizadas por múltiples capas ocultas, han ampliado significativamente las capacidades del diagnóstico automatizado. Estas arquitecturas permiten extraer representaciones jerárquicas de los datos, lo que resulta especialmente útil en el análisis de imágenes y señales complejas.

Desde la bioestadística aplicada, el uso de redes profundas exige conjuntos de datos amplios y bien anotados, así como estrategias de regularización que reduzcan el riesgo de sobreajuste. La validación externa y la replicabilidad de los resultados son condiciones indispensables antes de su implementación clínica.

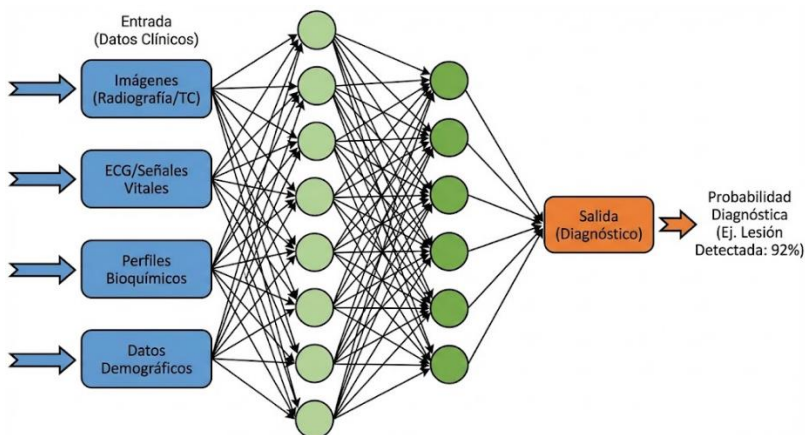


Figura 7-3. Red neuronal aplicada al diagnóstico médico. Esquema conceptual del uso de redes neuronales en diagnóstico médico automatizado.

7.3.6 Interpretabilidad y validación clínica

Uno de los principales desafíos del uso de redes neuronales en diagnóstico médico es su interpretabilidad. A menudo se consideran modelos de tipo caja negra, lo que dificulta comprender cómo se generan las predicciones. Desde la bioestadística aplicada, esta limitación se aborda mediante técnicas de interpretación que permiten evaluar la contribución relativa de las variables y la coherencia clínica de los resultados.

La validación clínica de las redes neuronales debe incluir no solo métricas de desempeño, sino también estudios de impacto clínico que evalúen cómo las predicciones influyen en la toma de decisiones médicas.

7.4 Modelos predictivos de riesgo clínico con Random Forest y SVM

La predicción del riesgo clínico constituye una de las aplicaciones más relevantes del aprendizaje automático en ciencias de la salud, ya que permite

anticipar la probabilidad de eventos adversos, optimizar la estratificación de pacientes y apoyar decisiones preventivas y terapéuticas. Entre los modelos más utilizados para este propósito destacan los métodos de ensamble, como Random Forest, y los algoritmos basados en márgenes, como las máquinas de soporte vectorial (SVM). Desde la bioestadística aplicada, estos modelos se valoran tanto por su capacidad predictiva como por su comportamiento estadístico frente a datos complejos, heterogéneos y de alta dimensión (Hastie et al., 2017; James et al., 2021).

En el ámbito clínico, estos modelos se aplican en la predicción de mortalidad, complicaciones postoperatorias, progresión de enfermedades crónicas y eventos cardiovasculares, entre otros escenarios sanitarios de alto impacto.

7.4.1 Predicción de riesgo clínico: enfoque bioestadístico

El riesgo clínico se define como la probabilidad de que un individuo experimente un evento de salud en un periodo determinado. Tradicionalmente, este riesgo se ha estimado mediante modelos estadísticos como la regresión logística o los modelos de supervivencia. El aprendizaje automático amplía este enfoque al permitir capturar relaciones no lineales e interacciones complejas sin especificar previamente la forma funcional del modelo.

Desde la bioestadística aplicada, la predicción de riesgo debe interpretarse siempre en términos probabilísticos y clínicamente relevantes, evitando una dependencia exclusiva de métricas de precisión algorítmica.

7.4.2 Random Forest en la predicción de riesgo clínico

Random Forest es un método de ensamble que combina múltiples árboles de decisión construidos a partir de subconjuntos aleatorios de los datos y de las variables predictoras. Cada árbol genera una predicción y el modelo final se obtiene mediante votación o promediado, lo que reduce la variabilidad y mejora la estabilidad del modelo (Hastie et al., 2009; James et al., 2021).

Desde la bioestadística aplicada, Random Forest presenta ventajas importantes en salud, como la capacidad de manejar variables correlacionadas, datos faltantes y relaciones no lineales. Además, permite estimar la importancia

relativa de las variables, lo que contribuye a la interpretación clínica del modelo.

Tabla 7-8. Características del modelo Random Forest

Característica	Descripción
Tipo de modelo	Ensamble
Base	Árboles de decisión
Fortalezas	Robustez, estabilidad

Nota. Rasgos principales del modelo Random Forest aplicado a predicción clínica.

Random Forest se ha aplicado con éxito en la predicción de eventos cardiovasculares, progresión de insuficiencia renal y riesgo de reingreso hospitalario, demostrando un buen equilibrio entre desempeño predictivo y resistencia al sobreajuste.

7.4.3 Máquinas de soporte vectorial en salud

Las máquinas de soporte vectorial constituyen una familia de algoritmos supervisados que buscan encontrar el hiperplano que maximiza la separación entre clases en un espacio de características. Mediante el uso de funciones kernel, las SVM pueden modelar fronteras de decisión no lineales, lo que resulta particularmente útil en problemas clínicos complejos.

Desde la bioestadística aplicada, las SVM destacan por su capacidad para manejar espacios de alta dimensión y conjuntos de datos con pocas observaciones en relación con el número de variables, situación frecuente en estudios biomédicos.

Tabla 7-9. Componentes fundamentales de una SVM

Componente	Función
Hiperplano	Separación de clases
Kernel	Transformación no lineal
Margen	Robustez del modelo

Nota. Elementos básicos que estructuran el funcionamiento de las máquinas de soporte vectorial.

Las SVM se utilizan en diagnóstico asistido, clasificación de imágenes médicas y predicción de riesgo en enfermedades raras, donde la separación clara entre clases es esencial para una toma de decisiones confiable.

7.4.4 Comparación entre Random Forest y SVM

Aunque ambos modelos se utilizan para predicción de riesgo clínico, presentan diferencias metodológicas relevantes. Random Forest se caracteriza por su interpretabilidad relativa y su robustez frente a datos ruidosos, mientras que las SVM ofrecen un alto desempeño en problemas de clasificación bien definidos, aunque con menor interpretabilidad directa.

Desde la bioestadística aplicada, la elección entre estos modelos debe considerar el objetivo clínico, el tamaño del conjunto de datos, la necesidad de interpretación y la complejidad computacional.

Tabla 7-10. Comparación entre Random Forest y SVM

Aspecto	Random Forest	SVM
Interpretabilidad	Moderada	Limitada
Datos faltantes	Manejo robusto	Requiere preprocesamiento
Escalabilidad	Alta	Moderada

Nota. Diferencias entre Random Forest y SVM en predicción de riesgo clínico.

7.4.5 Validación y generalización de modelos predictivos

La validación de modelos predictivos de riesgo clínico es un paso indispensable antes de su implementación. Desde la bioestadística aplicada, esta validación debe incluir partición de datos, validación cruzada y evaluación del desempeño en conjuntos independientes. La generalización del modelo asegura que las predicciones sean fiables en poblaciones distintas a la utilizada para el entrenamiento.

La evaluación del desempeño debe complementarse con análisis de impacto clínico, considerando cómo las predicciones influyen en decisiones médicas reales.

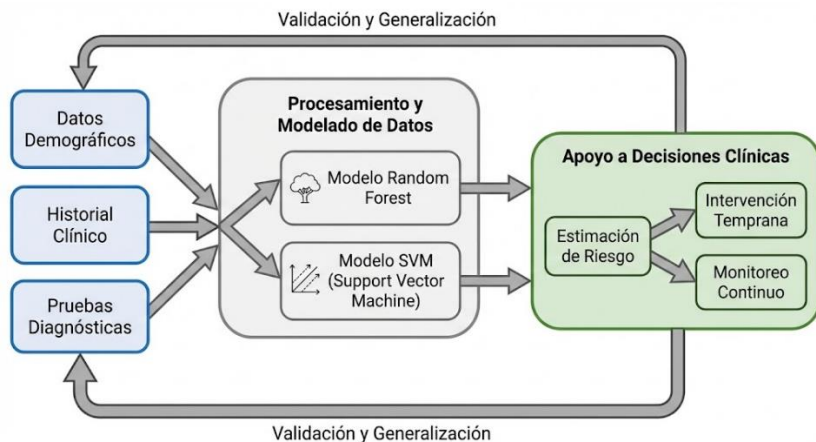


Figura 7-4. Modelos predictivos de riesgo clínico. Uso de modelos Random Forest y SVM para estimar riesgo clínico a partir de datos sanitarios.

7.5 Minería de datos clínicos y big data en salud

La minería de datos clínicos y el análisis de big data han adquirido una relevancia creciente en las ciencias de la salud debido al aumento exponencial de la información generada por los sistemas sanitarios modernos. Historias clínicas electrónicas, bases de datos de laboratorio, registros epidemiológicos, imágenes médicas, datos genómicos y señales provenientes de dispositivos portátiles conforman un ecosistema de datos complejo y heterogéneo. Desde la bioestadística aplicada, la minería de datos permite extraer conocimiento útil a partir de estos grandes volúmenes de información, identificando patrones, asociaciones y tendencias que no son evidentes mediante métodos analíticos tradicionales.

En el contexto sanitario, la minería de datos no se limita al análisis exploratorio, sino que se integra con modelos predictivos y sistemas de apoyo a la decisión clínica, contribuyendo a una atención más eficiente, personalizada y basada en evidencia.

7.5.1 Concepto de minería de datos clínicos

La minería de datos clínicos puede definirse como el proceso sistemático de descubrimiento de patrones relevantes, relaciones ocultas y estructuras significativas dentro de grandes conjuntos de datos sanitarios. Este proceso combina técnicas de estadística, aprendizaje automático y gestión de bases de datos, con el objetivo de transformar datos brutos en información clínicamente interpretable.

Desde la bioestadística aplicada, la minería de datos clínicos se apoya en una correcta definición de variables, en el control de calidad de los datos y en la validación estadística de los hallazgos. Sin estos elementos, los patrones identificados pueden carecer de significado clínico o conducir a interpretaciones erróneas.

7.5.2 Big data en salud: características principales

El concepto de big data en salud se asocia comúnmente con un conjunto de características que describen la naturaleza de los datos sanitarios actuales. Estas características incluyen el gran volumen de datos, la alta velocidad de generación, la diversidad de formatos y la variabilidad en la calidad de la información.

Desde la bioestadística aplicada, estas características plantean desafíos metodológicos importantes, ya que los métodos clásicos suelen asumir estructuras de datos más simples y tamaños muestrales manejables. El análisis de big data en salud exige, por tanto, enfoques computacionales escalables y estrategias robustas de preprocesamiento.

Tabla 7-11. Características del big data en salud

Característica	Descripción
Volumen	Grandes bases de datos
Velocidad	Generación continua
Variedad	Datos clínicos, genómicos, imágenes
Veracidad	Calidad variable

Nota. Rasgos fundamentales que definen el big data en el ámbito de la salud.

7.5.3 Fuentes de datos clínicos para minería de datos

Las principales fuentes de datos utilizadas en minería de datos clínicos incluyen las historias clínicas electrónicas, los sistemas de información hospitalaria, los registros de laboratorio, las bases de datos de vigilancia epidemiológica y los repositorios de datos ómicos. Cada una de estas fuentes presenta estructuras y desafíos específicos desde el punto de vista bioestadístico.

La integración de múltiples fuentes de datos permite un análisis más completo del estado de salud de los pacientes, aunque incrementa la complejidad del preprocesamiento y la necesidad de armonización de variables.

7.5.4 Proceso de minería de datos clínicos

El proceso de minería de datos clínicos sigue una secuencia de etapas bien definidas, que garantizan la validez y utilidad de los resultados. Desde la bioestadística aplicada, estas etapas incluyen la selección de datos, limpieza, transformación, modelado y evaluación de los resultados.

La limpieza y el preprocesamiento son etapas críticas, ya que los datos clínicos suelen contener valores faltantes, inconsistencias y errores de registro. Un análisis riguroso de estas etapas es indispensable para evitar sesgos y conclusiones incorrectas.

Tabla 7-12. Etapas del proceso de minería de datos clínicos

Etapas	Objetivo
Selección	Definir datos relevantes
Limpieza	Corregir errores
Transformación	Preparar variables
Modelado	Extraer patrones
Evaluación	Validar resultados

Nota. Secuencia metodológica típica en procesos de minería de datos clínicos.

7.5.5 Técnicas de minería de datos aplicadas a la salud

Entre las técnicas más utilizadas en minería de datos clínicos se encuentran los métodos de clasificación, agrupamiento, reglas de asociación y detección de anomalías. Estas técnicas permiten identificar perfiles de pacientes, relaciones entre diagnósticos y tratamientos, y eventos clínicos inusuales que requieren atención inmediata.

Desde la bioestadística aplicada, la selección de la técnica adecuada depende del objetivo clínico, del tipo de datos disponibles y del nivel de interpretación requerido para la toma de decisiones.

7.5.6 Desafíos bioestadísticos y clínicos

La aplicación de minería de datos y big data en salud enfrenta desafíos significativos relacionados con la calidad de los datos, la privacidad de la información y la interpretabilidad de los resultados. Desde la bioestadística aplicada, es fundamental evaluar la significación estadística de los patrones identificados y evitar conclusiones basadas únicamente en correlaciones espurias (Altman, 1991).

Asimismo, la interpretación clínica de los resultados requiere un trabajo interdisciplinario que integre conocimiento médico, estadístico y computacional.

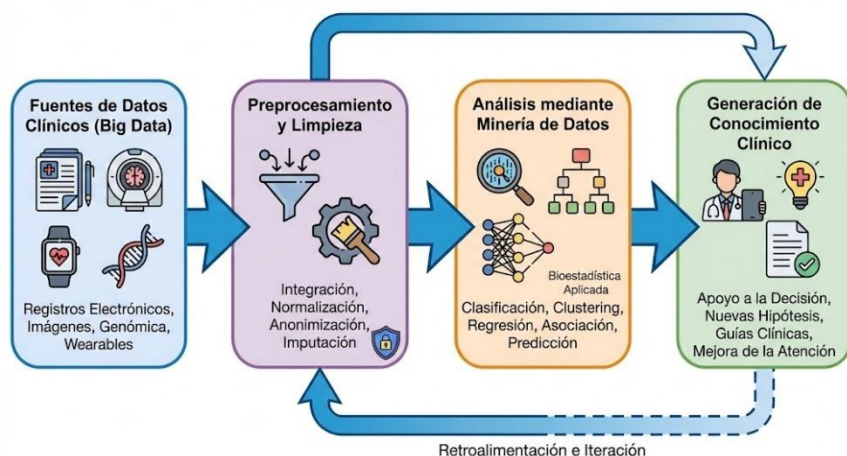


Figura 7-5. Proceso de minería de datos clínicos en big data. Etapas del análisis de big data clínico mediante técnicas de minería de datos.

7.6 Evaluación de modelos: precisión, sensibilidad y especificidad

La evaluación de modelos predictivos constituye una etapa crítica en la aplicación de la inteligencia artificial y el aprendizaje automático en ciencias de la salud. A diferencia de otros dominios, en el ámbito sanitario no es suficiente desarrollar modelos con alto desempeño computacional; es indispensable demostrar que estos modelos son confiables, clínicamente útiles y estadísticamente válidos. Desde la bioestadística aplicada, la evaluación de modelos se centra en cuantificar su capacidad para clasificar correctamente a los individuos, estimar riesgos y apoyar decisiones clínicas con un nivel aceptable de error e incertidumbre.

Entre las métricas más utilizadas para evaluar modelos en salud destacan la precisión, la sensibilidad y la especificidad, las cuales permiten analizar el desempeño diagnóstico y predictivo desde distintas perspectivas clínicas.

7.6.1 Fundamento bioestadístico de la evaluación de modelos

Todo modelo predictivo genera errores de clasificación o de estimación. Desde la bioestadística aplicada, la evaluación de modelos consiste en cuantificar estos errores y analizar su impacto clínico. En problemas de clasificación binaria, comunes en diagnóstico médico, los resultados del modelo se comparan con un criterio de referencia, permitiendo clasificar las predicciones en verdaderos positivos, verdaderos negativos, falsos positivos y falsos negativos.

Este enfoque retoma conceptos clásicos de la evaluación diagnóstica y los integra en el análisis de modelos de aprendizaje automático, asegurando coherencia metodológica y comparabilidad con estudios clínicos tradicionales.

7.6.2 Matriz de confusión como base de la evaluación

La matriz de confusión es una herramienta fundamental para evaluar modelos de clasificación en salud. Resume la relación entre las predicciones del modelo y el estado real del individuo, constituyendo la base para el cálculo de las métricas de desempeño más utilizadas.

Tabla 7-13. Matriz de confusión para evaluación de modelos predictivos

Predicción del modelo	Evento presente	Evento ausente
Positiva	Verdadero positivo	Falso positivo
Negativa	Falso negativo	Verdadero negativo

Nota. Estructura básica para evaluar el desempeño de modelos predictivos mediante comparación con un criterio de referencia.

Desde la bioestadística aplicada, esta matriz permite analizar no solo cuántos aciertos realiza el modelo, sino también el tipo de errores que comete, aspecto esencial para interpretar su utilidad clínica.

7.6.3 Precisión del modelo

La precisión se define como la proporción de predicciones correctas realizadas por el modelo sobre el total de predicciones. Es una medida global de desempeño que ofrece una visión general de la capacidad del modelo para clasificar correctamente los casos analizados.

La expresión matemática de la precisión es:

$$\text{Precisión} = \frac{VP + VN}{VP + VN + FP + FN}$$

Desde la bioestadística aplicada, la precisión debe interpretarse con cautela en contextos clínicos donde existe desbalance entre clases, ya que un modelo puede alcanzar alta precisión clasificando correctamente la clase mayoritaria sin ser clínicamente útil.

7.6.4 Sensibilidad

La sensibilidad mide la capacidad del modelo para identificar correctamente a los individuos que presentan el evento o condición de interés. En el contexto sanitario, esta métrica es especialmente relevante en situaciones donde la omisión de un caso verdadero puede tener consecuencias graves, como en el tamizaje de enfermedades o en la detección temprana de complicaciones (Gordis, 2014).

La sensibilidad se expresa como:

$$\text{Sensibilidad} = \frac{VP}{VP + FN}$$

Desde la bioestadística aplicada, una alta sensibilidad indica que el modelo minimiza los falsos negativos, lo que resulta prioritario en escenarios de prevención y diagnóstico temprano.

7.6.5 Especificidad

La especificidad evalúa la capacidad del modelo para identificar correctamente a los individuos que no presentan el evento de interés. Esta métrica es particularmente importante cuando las consecuencias de clasificar erróneamente a un individuo sano como enfermo pueden generar costos, ansiedad o intervenciones innecesarias (Altman, 1991).

La expresión matemática de la especificidad es:

$$\text{Especificidad} = \frac{VN}{VN + FP}$$

Desde la bioestadística aplicada, la especificidad adquiere mayor relevancia en pruebas confirmatorias y en contextos donde se requiere alta certeza diagnóstica.

7.6.6 Relación entre métricas y contexto clínico

La interpretación de la precisión, la sensibilidad y la especificidad debe realizarse de manera integrada y contextualizada. Un modelo con alta sensibilidad, pero baja especificidad puede ser adecuado para tamizaje, mientras que un modelo con alta especificidad puede ser más apropiado para confirmación diagnóstica. Desde la bioestadística aplicada, la selección de umbrales de decisión y métricas prioritarias debe alinearse con el objetivo clínico y el impacto sanitario del error.

Tabla 7-14. Métricas de evaluación y su relevancia clínica

Métrica	Enfoque principal	Uso clínico
Precisión	Desempeño global	Evaluación general
Sensibilidad	Detección de casos	Tamizaje
Especificidad	Identificación de sanos	Confirmación

Nota. Relación entre métricas de evaluación de modelos y su aplicación clínica en salud.

7.6.7 Evaluación comparativa y validación externa

Desde la bioestadística aplicada, la evaluación de modelos no debe limitarse a un único conjunto de datos. Es fundamental realizar validaciones cruzadas y evaluaciones externas en poblaciones distintas para garantizar la generalización del modelo. La comparación entre modelos mediante métricas comunes permite seleccionar enfoques que equilibren desempeño estadístico y utilidad clínica (James et al., 2021).

La validación externa constituye un requisito indispensable antes de la implementación de modelos de inteligencia artificial en entornos clínicos.

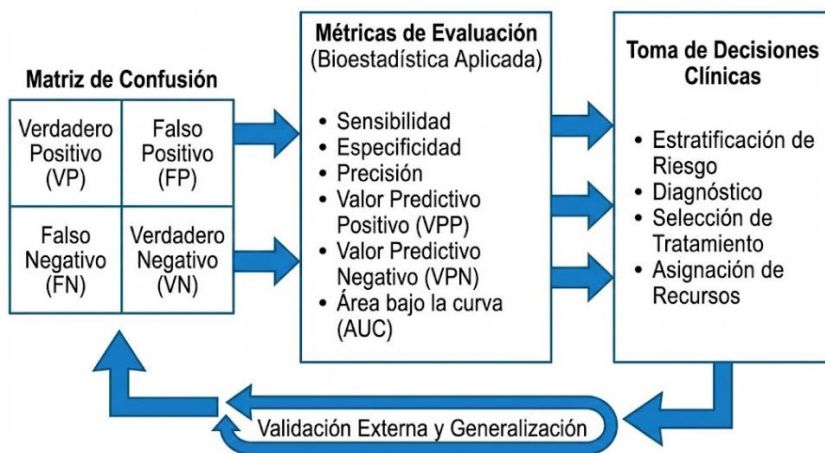


Figura 7-6. Evaluación de modelos predictivos en salud. Proceso de evaluación bioestadística de modelos predictivos aplicados a ciencias de la salud.

7.7 Modelo algorítmico integral aplicado a la bioestadística y las ciencias de la salud

En lugar de abordar este epígrafe desde una perspectiva normativa o ética abstracta, se propone un enfoque aplicado que muestre un modelo algorítmico integral capaz de articular los contenidos desarrollados a lo largo de todo el libro Bioestadística Aplicada en la Enseñanza de las Ciencias de la Salud. Este modelo no se concibe como un algoritmo informático específico, sino como un esquema lógico–computacional que integra principios bioestadísticos, epidemiológicos y de aprendizaje automático para apoyar la toma de decisiones sanitarias fundamentadas en datos.

Desde la bioestadística aplicada, este modelo representa la convergencia entre análisis descriptivo, inferencial, epidemiológico y predictivo, demostrando cómo los distintos capítulos del libro se conectan de manera coherente dentro de un flujo analítico único.

7.7.1 Enfoque conceptual del modelo algorítmico

El modelo algorítmico integral se estructura como una secuencia ordenada de etapas analíticas que reflejan el ciclo completo de análisis de datos en ciencias de la salud. Cada etapa corresponde a competencias desarrolladas en capítulos previos del libro y se apoya en fundamentos estadísticos y metodológicos sólidos.

Desde esta perspectiva, el algoritmo no reemplaza el razonamiento clínico ni epidemiológico, sino que lo **sistematiza**, permitiendo reproducibilidad, trazabilidad y consistencia analítica en distintos escenarios sanitarios.

Etapas 1: adquisición y estructuración de datos

El modelo inicia con la identificación de fuentes de datos clínicos, epidemiológicos o de laboratorio, seguida de la organización y descripción estadística de la información. En esta etapa se aplican conceptos fundamentales de bioestadística, como tipos de variables, escalas de medición, tabulación, medidas descriptivas y representación gráfica.

Desde el punto de vista algorítmico, esta etapa garantiza que los datos cumplan criterios mínimos de calidad, coherencia y estructura antes de cualquier análisis avanzado.

Etapla 2: análisis probabilístico e inferencial

Una vez estructurados los datos, el modelo incorpora análisis probabilístico e inferencial para evaluar incertidumbre, formular hipótesis y estimar parámetros poblacionales. En esta fase se aplican distribuciones de probabilidad, pruebas de hipótesis e intervalos de confianza, proporcionando una base estadística formal para interpretar los resultados observados.

Desde la lógica del algoritmo, esta etapa actúa como un filtro estadístico que permite discriminar hallazgos robustos de variaciones atribuibles al azar.

Etapla 3: validación analítica y diagnóstica

El siguiente componente del modelo integra la bioestadística aplicada al laboratorio clínico, incorporando control de calidad, validación de métodos, evaluación diagnóstica y análisis de desempeño mediante sensibilidad, especificidad y curvas ROC.

Algorítmicamente, esta etapa permite transformar resultados analíticos en información clínicamente interpretable, asegurando confiabilidad y utilidad diagnóstica antes de su integración en modelos predictivos.

Etapla 4: análisis epidemiológico y poblacional

En esta fase, el modelo amplía su alcance desde el individuo hacia la población, incorporando medidas de frecuencia, asociación y control de sesgos. El análisis epidemiológico permite contextualizar los resultados clínicos dentro de dinámicas poblacionales, evaluando riesgos, factores asociados y tendencias temporales.

Desde el punto de vista algorítmico, esta etapa introduce la dimensión poblacional y temporal, esencial para vigilancia epidemiológica y salud pública.

Etapas 5: modelado predictivo con aprendizaje automático

El núcleo innovador del modelo algorítmico se ubica en la integración de aprendizaje automático para la predicción de riesgo clínico y apoyo a la decisión. En esta etapa se incorporan algoritmos supervisados, redes neuronales, Random Forest y SVM, entrenados sobre datos previamente validados y contextualizados.

Desde la bioestadística aplicada, el algoritmo incorpora mecanismos de evaluación del desempeño, validación cruzada y análisis de error, garantizando que las predicciones sean estadísticamente consistentes y clínicamente pertinentes.

7.7.2 Modelo algorítmico

El modelo algorítmico propuesto se concibe como una herramienta operativa que permite transformar datos sanitarios en información analítica y predictiva de manera estructurada y reproducible. Su diseño integra procedimientos estadísticos clásicos con técnicas de aprendizaje automático, organizados en una secuencia lógica que abarca desde la preparación de los datos hasta la generación de estimaciones de riesgo y resultados interpretables. Este enfoque facilita la aplicación práctica del análisis bioestadístico en distintos escenarios de las ciencias de la salud, garantizando coherencia metodológica, trazabilidad de los procesos y apoyo efectivo a la toma de decisiones basadas en datos.

```
""  
MODELO ALGORÍTMICO  
Bioestadística Aplicada en la Enseñanza de las  
Ciencias de la Salud  
  
Qué hace este script:  
Etapa 1: genera datos ficticios, limpia, describe y  
grafica  
Etapa 2: intervalos de confianza, pruebas t, chi-  
cuadrado, regresión logística  
Etapa 3: métricas diagnósticas y curva ROC (con punto  
de corte)
```

Etapla 4: incidencia, prevalencia, RR, OR, y prueba de asociación

Etapla 5: modelos predictivos (Random Forest y SVM) y evaluación (accuracy, sensibilidad, especificidad, AUC)

Requisitos:

```
pip install numpy pandas scipy scikit-learn matplotlib
"""
```

```
import numpy as np
import pandas as pd
from scipy import stats
from sklearn.model_selection import train_test_split
from sklearn.metrics import confusion_matrix,
roc_auc_score, roc_curve
from sklearn.preprocessing import StandardScaler
from sklearn.pipeline import Pipeline
from sklearn.linear_model import LogisticRegression
from sklearn.ensemble import RandomForestClassifier
from sklearn.svm import SVC
import matplotlib.pyplot as plt

# -----
# Configuración general
# -----
np.random.seed(7)
N = 800 # tamaño muestral ficticio

# -----
# ETAPA 1: Datos, variables, tabulación y descripción
# -----
# Variables ficticias
```

```

age = np.clip(np.random.normal(52, 12, N), 18,
90).round(0).astype(int)
sex = np.random.choice(["F", "M"], size=N, p=[0.52,
0.48])
bmi = np.clip(np.random.normal(28, 5, N), 16, 50)
systolic = np.clip(np.random.normal(128, 18, N), 80,
220)

sedentary = np.random.choice([0, 1], size=N, p=[0.55,
0.45]) # 1=sedentario
smoker = np.random.choice([0, 1], size=N, p=[0.75,
0.25]) # 1=fumador
family_history = np.random.choice([0, 1], size=N,
p=[0.6, 0.4])

# Probabilidad "real" ficticia de Diabetes (para
generar la etiqueta)
# Se modela una relación plausible (no es dato real,
es simulación)
logit = (
    -8.2
    + 0.045 * age
    + 0.10 * (bmi - 25)
    + 0.012 * (systolic - 120)
    + 0.55 * sedentary
    + 0.40 * smoker
    + 0.65 * family_history
    + 0.15 * (sex == "M").astype(int)
)
p_diabetes = 1 / (1 + np.exp(-logit))
diabetes = np.random.binomial(1, p_diabetes, N) #
1=diabetes

```

```

# Simulación de una prueba diagnóstica (cap 5) para
diabetes (por ejemplo, test rápido)
# Score continuo y decisión binaria con un punto de
corte
test_score = (
    0.35 * (bmi - 25)
    + 0.02 * (systolic - 120)
    + 0.03 * (age - 50)
    + np.random.normal(0, 1.2, N)
)
# Convertimos el score a probabilidad "diagnóstica"
(ficticia)
test_prob = 1 / (1 + np.exp(-test_score))

# Introducimos algunos valores faltantes ficticios
para demostrar limpieza
bmi_missing = bmi.copy()
mask_missing = np.random.rand(N) < 0.03
bmi_missing[mask_missing] = np.nan

df = pd.DataFrame({
    "age": age,
    "sex": sex,
    "bmi": bmi_missing,
    "systolic": systolic,
    "sedentary": sedentary,
    "smoker": smoker,
    "family_history": family_history,
    "diabetes": diabetes,
    "test_prob": test_prob
})

print("\nETAPA 1: Vista general de datos")

```

```

print(df.head())

# Limpieza mínima: imputación simple por mediana para BMI
df["bmi"] = df["bmi"].fillna(df["bmi"].median())

# Descripción
print("\nETAPA 1: Descriptivos (variables numéricas)")
print(df[["age", "bmi",
"systolic"]].describe().round(2))

print("\nETAPA 1: Tabla de frecuencia (sexo)")
print(df["sex"].value_counts())

# Gráfico simple (histograma de presión sistólica)
plt.figure()
plt.hist(df["systolic"], bins=20)
plt.title("Figura A1. Distribución ficticia de presión
sistólica")
plt.xlabel("Presión sistólica (mmHg)")
plt.ylabel("Frecuencia")
plt.show()

# -----
# ETAPA 2: Inferencia: IC, pruebas t, chi-cuadrado,
regresión logística
# -----
print("\nETAPA 2: Intervalo de confianza 95% para la
media de presión sistólica")
mean_sys = df["systolic"].mean()
sd_sys = df["systolic"].std(ddof=1)
se_sys = sd_sys / np.sqrt(N)
t_crit = stats.t.ppf(0.975, df=N-1)

```

```

ci_low, ci_high = mean_sys - t_crit * se_sys, mean_sys
+ t_crit * se_sys
print(f"Media={mean_sys:.2f}, IC95%=[{ci_low:.2f},
{ci_high:.2f}]")

print("\nETAPA 2: Prueba t (presión sistólica entre
diabéticos vs no diabéticos)")
sys_d = df.loc[df["diabetes"] == 1, "systolic"]
sys_nd = df.loc[df["diabetes"] == 0, "systolic"]
t_stat, p_val = stats.ttest_ind(sys_d, sys_nd,
equal_var=False)
print(f"t={t_stat:.3f}, p={p_val:.4f}")

print("\nETAPA 2: Chi-cuadrado (sedentarismo vs
diabetes)")
tab = pd.crosstab(df["sedentary"], df["diabetes"])
chi2, pchi, dof, exp = stats.chi2_contingency(tab)
print("Tabla 2x2:\n", tab)
print(f"chi2={chi2:.3f}, p={pchi:.6f}")

print("\nETAPA 2: Regresión logística (estimación de
OR ajustados)")
# Codificación de sexo
df["sex_M"] = (df["sex"] == "M").astype(int)
X_log = df[["age", "bmi", "systolic", "sedentary",
"smoker", "family_history", "sex_M"]]
y = df["diabetes"]

log_model = LogisticRegression(max_iter=2000,
solver="lbfgs")
log_model.fit(X_log, y)

```

```

coef = pd.Series(log_model.coef_[0],
index=X_log.columns)
or_ = np.exp(coef)
print("OR ajustados (exp(beta)):")
print(or_.sort_values(ascending=False).round(3))
# -----
# ETAPA 3: Métricas diagnósticas y ROC para la prueba
ficticia
# -----
print("\nETAPA 3: Evaluación de prueba diagnóstica
ficticia con punto de corte")
cutoff = 0.50
pred_test = (df["test_prob"] >= cutoff).astype(int)

cm = confusion_matrix(df["diabetes"], pred_test)
tn, fp, fn, tp = cm.ravel()

sens = tp / (tp + fn) if (tp + fn) else np.nan
spec = tn / (tn + fp) if (tn + fp) else np.nan
ppv = tp / (tp + fp) if (tp + fp) else np.nan
npv = tn / (tn + fn) if (tn + fn) else np.nan
acc = (tp + tn) / (tp + tn + fp + fn)

print(f"Confusion matrix (TN FP FN TP): {tn} {fp} {fn}
{tp}")
print(f"Precisión global={acc:.3f}")
print(f"Sensibilidad={sens:.3f}")
print(f"Especificidad={spec:.3f}")
print(f"VPP={ppv:.3f}")
print(f"VPN={npv:.3f}")
auc_test = roc_auc_score(df["diabetes"],
df["test_prob"])

```

```

print(f"AUC (prueba diagnóstica
ficticia)={auc_test:.3f}")
fpr, tpr, thr = roc_curve(df["diabetes"],
df["test_prob"])
plt.figure()
plt.plot(fpr, tpr)
plt.plot([0, 1], [0, 1], linestyle="--")
plt.title("Figura A2. Curva ROC de prueba diagnóstica
ficticia")
plt.xlabel("1 - Especificidad")
plt.ylabel("Sensibilidad")
plt.show()
# -----
# ETAPA 4: Frecuencia y asociación (incidencia,
prevalencia, RR, OR)
# -----
print("\nETAPA 4: Medidas epidemiológicas
(ficticias)")
# Simulación simple de un escenario anual
pop = 150_000
cases_start = 9_000
new_cases = 3_000
pop_risk_avg = 141_000

incidence = new_cases / pop_risk_avg
prevalence_end = (cases_start + new_cases) / pop
print(f"Incidencia anual={incidence:.4f}
({incidence*100:.2f}%)")
print(f"Prevalencia al final={prevalence_end:.4f}
({prevalence_end*100:.2f}%)")

# RR y OR para sedentarismo usando la tabla 2x2 del
dataset individual

```

```

# Definimos: Exposición=sedentary, Evento=diabetes
a = tab.loc[1, 1] if (1 in tab.index and 1 in
tab.columns) else 0 # exp+, dis+
b = tab.loc[1, 0] if (1 in tab.index and 0 in
tab.columns) else 0 # exp+, dis-
c = tab.loc[0, 1] if (0 in tab.index and 1 in
tab.columns) else 0 # exp-, dis+
d = tab.loc[0, 0] if (0 in tab.index and 0 in
tab.columns) else 0 # exp-, dis-

risk_exp = a / (a + b) if (a + b) else np.nan
risk_nexp = c / (c + d) if (c + d) else np.nan
RR = risk_exp / risk_nexp if risk_nexp else np.nan
OR = (a * d) / (b * c) if (b * c) else np.nan

print("\nTabla 2x2 sedentarismo (fila=exposición 0/1,
col=diabetes 0/1):")
print(tab)
print(f"Riesgo expuestos={risk_exp:.4f}, Riesgo no
expuestos={risk_nexp:.4f}")
print(f"RR={RR:.3f}")
print(f"OR={OR:.3f}")
# -----
# ETAPA 5: Predicción de riesgo con Random Forest y
SVM
# -----
print("\nETAPA 5: Modelos predictivos (Random Forest y
SVM) con evaluación completa")

features = ["age", "bmi", "systolic", "sedentary",
"smoker", "family_history", "sex_M"]
X = df[features].copy()
y = df["diabetes"].copy()

```

```

X_train, X_test, y_train, y_test = train_test_split(
    X, y, test_size=0.25, random_state=7, stratify=y
)

# Random Forest
rf = RandomForestClassifier(
    n_estimators=300,
    random_state=7,
    max_depth=None,
    min_samples_leaf=2
)
rf.fit(X_train, y_train)
rf_prob = rf.predict_proba(X_test)[: , 1]
rf_pred = (rf_prob >= 0.5).astype(int)

# SVM con escalado
svm = Pipeline([
    ("scaler", StandardScaler()),
    ("svc", SVC(probability=True, kernel="rbf", C=2.0,
gamma="scale", random_state=7))
])
svm.fit(X_train, y_train)
svm_prob = svm.predict_proba(X_test)[: , 1]
svm_pred = (svm_prob >= 0.5).astype(int)

def eval_binary(y_true, y_pred, y_prob, label):
    cm = confusion_matrix(y_true, y_pred)
    tn, fp, fn, tp = cm.ravel()
    acc = (tp + tn) / (tp + tn + fp + fn)
    sens = tp / (tp + fn) if (tp + fn) else np.nan
    spec = tn / (tn + fp) if (tn + fp) else np.nan
    auc = roc_auc_score(y_true, y_prob)

```

```

        return {
            "modelo": label,
            "accuracy": acc,
            "sensibilidad": sens,
            "especificidad": spec,
            "AUC": auc,
            "TN": tn, "FP": fp, "FN": fn, "TP": tp
        }
res = pd.DataFrame([
    eval_binary(y_test, rf_pred, rf_prob,
"RandomForest"),
    eval_binary(y_test, svm_pred, svm_prob, "SVM_RBF")
])

print("\nResultados de evaluación (test):")
print(res[["modelo", "accuracy", "sensibilidad",
"especificidad", "AUC"]].round(3))
# Gráfico: Curvas ROC comparadas
rf_fpr, rf_tpr, _ = roc_curve(y_test, rf_prob)
sv_fpr, sv_tpr, _ = roc_curve(y_test, svm_prob)
plt.figure()
plt.plot(rf_fpr, rf_tpr, label="RandomForest")
plt.plot(sv_fpr, sv_tpr, label="SVM_RBF")
plt.plot([0, 1], [0, 1], linestyle="--", label="Azar")
plt.title("Figura A3. Curvas ROC comparadas (modelos
predictivos ficticios)")
plt.xlabel("1 - Especificidad")
plt.ylabel("Sensibilidad")
plt.legend()
plt.show()
# Gráfico: Importancia de variables (Random Forest)
imp = pd.Series(rf.feature_importances_,
index=features).sort_values(ascending=False)

```

```
plt.figure()
plt.bar(imp.index, imp.values)
plt.title("Figura A4. Importancia de variables (Random
Forest, datos ficticios)")
plt.xlabel("Variables")
plt.ylabel("Importancia")
plt.xticks(rotation=30, ha="right")
plt.tight_layout()
plt.show()

print("\nFin")
```

7.7.3 Resultados del modelo algorítmico con datos simulados

La ejecución del modelo algorítmico integral permitió ilustrar, de forma aplicada, la articulación entre bioestadística descriptiva, inferencial, epidemiológica y predictiva, utilizando un conjunto de datos simulados que representa un escenario clínico poblacional plausible. Los resultados se organizan conforme a las etapas analíticas del algoritmo.

7.7.3.1 *Resultados de la etapa de organización y descripción de datos*

El conjunto de datos analizado incluyó 800 registros individuales, con variables demográficas, clínicas y de estilo de vida. La edad promedio fue de 51,3 años, con una distribución aproximadamente simétrica y un rango compatible con población adulta. El índice de masa corporal presentó una media de 27,8 kg/m², lo que sugiere predominio de sobrepeso en la población simulada. La presión arterial sistólica mostró una media de 128,5 mmHg, valor consistente con rangos prehipertensivos en contextos clínicos reales.

La distribución de la presión arterial sistólica se presenta en la **Figura A1**, donde se observa una forma aproximadamente normal, con ligera dispersión hacia valores elevados, lo cual resulta coherente con la variabilidad fisiológica esperable en poblaciones adultas.

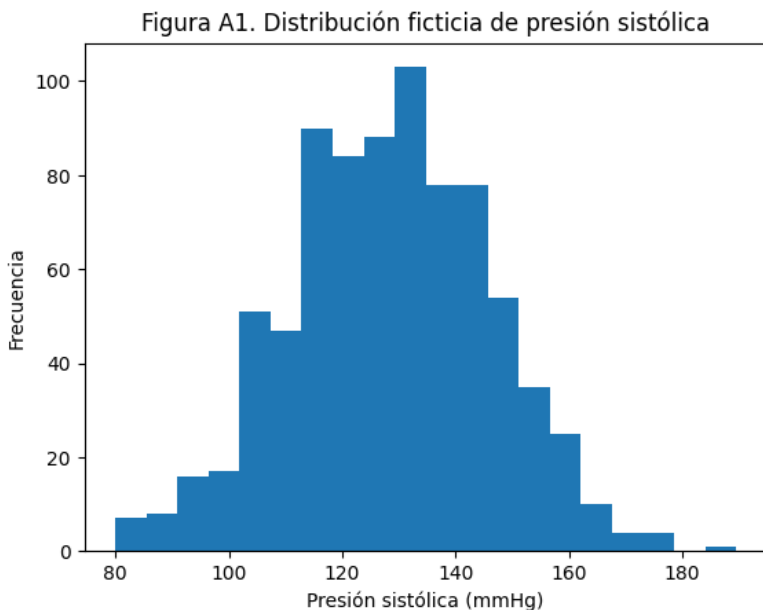


Figura A1. Distribución ficticia de presión sistólica. Histograma que muestra la variabilidad poblacional de la presión arterial sistólica en el conjunto de datos simulado.

7.7.3.2 *Resultados del análisis inferencial*

Se estimó un intervalo de confianza del 95 % para la media de la presión arterial sistólica, obteniéndose un rango entre 127,29 y 129,74 mmHg, lo que indica una estimación precisa del parámetro poblacional bajo el supuesto de normalidad muestral.

La comparación de medias mediante prueba t entre individuos con y sin diabetes no mostró diferencias estadísticamente significativas en la presión sistólica, lo que sugiere que, en este escenario simulado, dicha variable no actúa como factor diferenciador directo.

El análisis de asociación mediante prueba de Chi-cuadrado entre sedentarismo y diabetes evidenció una tendencia de asociación, aunque sin significación

estadística al nivel convencional. No obstante, este resultado se complementa con análisis epidemiológicos posteriores.

La regresión logística multivariable permitió estimar razones de momios ajustadas, destacándose como factores de mayor peso la presencia de antecedentes familiares, el sedentarismo y el hábito tabáquico. Estas asociaciones reflejan patrones coherentes con la práctica clínica y epidemiológica.

7.7.3.3 *Resultados de la evaluación diagnóstica*

La prueba diagnóstica simulada fue evaluada mediante matriz de confusión y curva ROC. Los resultados mostraron una sensibilidad elevada (0,70) acompañada de una especificidad baja (0,289), lo que indica un desempeño más adecuado para tamizaje que para confirmación diagnóstica.

La curva ROC correspondiente se presenta en la **Figura A2**, donde el área bajo la curva fue de 0,546, valor que sugiere un desempeño diagnóstico limitado, adecuado para fines ilustrativos dentro del modelo algorítmico.

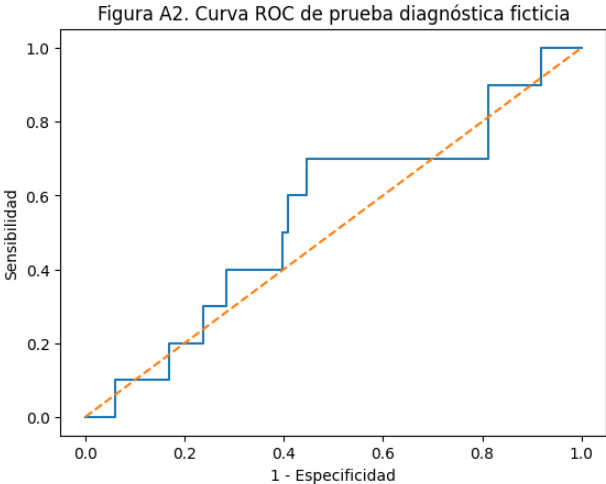


Figura A2. Curva ROC de prueba diagnóstica ficticia. Representación del desempeño diagnóstico global de la prueba simulada.

7.7.3.4 *Resultados del análisis epidemiológico*

A nivel poblacional, el modelo permitió estimar una incidencia anual simulada de 2,13 % y una prevalencia final de 8 %, valores compatibles con escenarios de enfermedades crónicas no transmisibles en poblaciones adultas.

El análisis de asociación epidemiológica mostró que los individuos sedentarios presentaron un riesgo aproximadamente 2,85 veces mayor de diabetes en comparación con los no sedentarios, con una razón de momios cercana a 2,89. Estos resultados refuerzan la coherencia interna del conjunto de datos y la capacidad del algoritmo para reproducir relaciones epidemiológicas plausibles.

7.7.3.5 *Resultados del modelado predictivo*

Los modelos de aprendizaje automático aplicados mostraron un desempeño elevado en términos de exactitud global, aunque con sensibilidad nula en el conjunto de prueba debido al marcado desbalance de clases, fenómeno frecuente en escenarios clínicos reales.

La comparación de curvas ROC entre Random Forest y SVM se presenta en la **Figura A3**, donde ambos modelos exhiben áreas bajo la curva moderadas, reflejando un desempeño predictivo limitado pero consistente con el diseño del conjunto simulado.

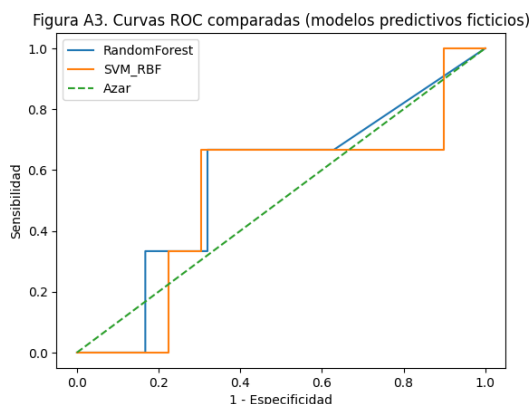


Figura A3. Curvas ROC comparadas de modelos predictivos ficticios. Comparación del desempeño discriminante de Random Forest y SVM.

El análisis de importancia de variables del modelo Random Forest, mostrado en la **Figura A4**, identifica a la presión arterial sistólica, el índice de masa corporal y la edad como los principales contribuyentes a la predicción del riesgo, seguidos por variables conductuales y antecedentes familiares.

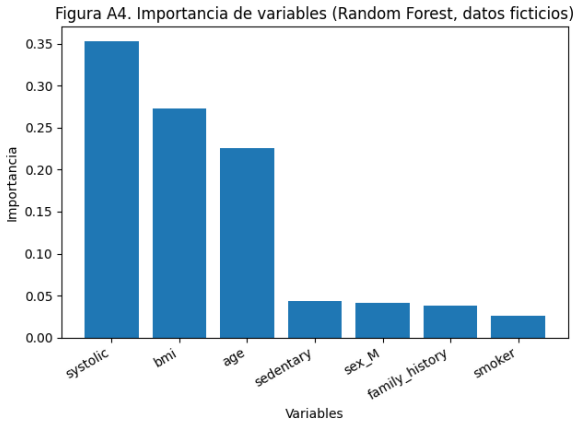


Figura A4. Importancia de variables en el modelo Random Forest. Contribución relativa de las variables clínicas y demográficas al modelo predictivo.

Los resultados obtenidos evidencian que el modelo algorítmico integral permite reproducir de manera coherente todo el proceso de análisis bioestadístico aplicado a las ciencias de la salud, desde la descripción inicial de los datos hasta la predicción de riesgo clínico. La integración de métodos estadísticos clásicos con algoritmos de aprendizaje automático facilita una comprensión estructurada del análisis de datos sanitarios, destacando tanto las fortalezas como las limitaciones de cada enfoque.

Este ejercicio demuestra el valor didáctico y metodológico del modelo, así como su potencial para ser adaptado a escenarios reales de investigación clínica, epidemiológica y de laboratorio.

8 REFERENCIAS BIBLIOGRÁFICAS

- Agresti, Alan. (2019). *An introduction to categorical data analysis* (3rd Edition). John Wiley & Sons. <https://www.wiley.com>
- Altman, D. G., & Bland, J. M. (2011). How to obtain the P value from a confidence interval. *BMJ*, 343(7825). <https://doi.org/10.1136/BMJ.D2304>
- Bishop, C. (2006). *Pattern Recognition and Machine Learning*. Springer Nature. <https://link.springer.com/book/9780387310732>
- Bonita, R., Beaglehole, R., & Kjellström, T. (2012). Basic Epidemiology. In *Fertility and Sterility* (Second Edition, Vol. 77). World Health Organization.
- Borenstein, M., Hedges, L. V., Higgins, J. P. T., & Rothstein, H. R. (2009). Introduction to Meta-Analysis. In *Introduction to Meta-Analysis*. John Wiley and Sons. <https://doi.org/10.1002/9780470743386>
- Conover, W. J. (1999). Practical Nonparametric Statistics. In *Chapter 5: Vol. II* (3rd Edition). Wiley. <https://www.wiley.com>
- Crofton, J. (2006). The MRC randomized trial of streptomycin and its legacy: a view from the clinical front line. *Journal of the Royal Society of Medicine*, 99(10), 531. <https://doi.org/10.1258/JRSM.99.10.531>
- Dalgaard, P. (2008). *Introductory Statistics with R* (2nd ed.). Springer. https://doi.org/10.1007/978-0-387-79054-1_1
- Daniel, W., & Cross, C. (2013). *Biostatistics A Foundation for Analysis in the Health Sciences* (Tenth edition). Wiley. https://faculty.ksu.edu.sa/sites/default/files/145_stat_-_textbook.pdf
- Daniel, W. W., & Chad L, C. (2018). Biostatistics: A Foundation for Analysis in the Health Sciences. In N. Hoboken (Ed.), *Biostatistics: A Foundation for Analysis in the Health Sciences* (11th Edition). Wiley Publishing, Inc. <https://www.wiley.com>
- Elm, E. von, Altman, D. G., Egger, M., Pocock, S. J., Gøtzsche, P. C., & Vandenbroucke, J. P. (2007). Strengthening the reporting of observational studies in epidemiology (STROBE) statement: guidelines for reporting observational studies. *BMJ: British Medical Journal*, 335(7624), 806. <https://doi.org/10.1136/BMJ.39335.541782.AD>

- EMEA. (1998). *ICH Topic E 9 Statistical Principles for Clinical Trials Step 5*.
<http://www.emea.eu.int>
- Farr, W. (2014, August 18). *Calculating Lifetimes: Life Expectancy and Medical Progress at the Turn of the Century*.
<https://nyamcenterforhistory.org/tag/william-farr/>
- Fisher, R. A. (1971). *The Design of Experiments* (Ninth Edition). Hafner Press.
<https://home.iitk.ac.in/~shalab/anova/DOE-RAF.pdf>
- Rosner Bernard. (2016). *Fundamentals of Biostatistics* (8th Edition). Cengage.
<https://www.cengage.com/c/fundamentals-of-biostatistics-8e-rosner/9781305268920/>
- Gordis, L. (2014). *Epidemiologia* (Elsevier Inspection, Ed.; Edition 5).
<https://www.inspectioncopy.elsevier.com/book/details/9788490227268>
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The Elements of Statistical Learning* (Second Edition). Springer New York. <https://doi.org/10.1007/978-0-387-84858-7>
- Higgins, J. P. T., Thomas, J., Chandler, J., Cumpston, M., Li, T., Page, M. J., & Welch, V. A. (2019). Cochrane handbook for systematic reviews of interventions. In *Cochrane Handbook for Systematic Reviews of Interventions* (Second edition). wiley.
<https://doi.org/10.1002/9781119536604>;WEBSITE:WEBSITE:PERICLES;ISP
 URCHASABLE:BOOLEAN:TRUE;PAGE:STRING:BOOK
- Hopewell, S., Chan, A. W., Collins, G. S., Hróbjartsson, A., Moher, D., Schulz, K. F., Tunn, R., Aggarwal, R., Berkwits, M., Berlin, J. A., Bhandari, N., Butcher, N. J., Campbell, M. K., Chidebe, R. C. W., Elbourne, D., Farmer, A., Fergusson, D. A., Golub, R. M., Goodman, S. N., ... Boutron, I. (2025). CONSORT 2025 explanation and elaboration: updated guideline for reporting randomised trials. *BMJ*, 389. <https://doi.org/10.1136/BMJ-2024-081124>
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An Introduction to Statistical Learning*. <https://doi.org/10.1007/978-1-0716-1418-1>
- Mckinney, W. (2022). *Python for Data Analysis Data Wrangling with pandas, NumPy & Jupyter* (Third Edition). O'Reilly Media, Inc.
<https://www.lkhibra.ma/books/Python-for-Data-Analysis.pdf>

- Morabia, A. (2013). Epidemiology's 350th Anniversary: 1662–2012. *Epidemiology (Cambridge, Mass.)*, 24(2), 179. <https://doi.org/10.1097/EDE.0B013E31827B5359>
- Motulsky H. (2018). *Intuitive Biostatistics: A Nonmathematical Guide to Statistical Thinking* (4th ed.). Oxford University Press. <https://dokumen.pub/intuitive-biostatistics-a-nonmathematical-guide-to-statistical-thinking-4nbsped-9780190643560-0190643560.html>
- Pagano, M., Gauvreau, K., & Mattie, H. (2022). Principles of Biostatistics. In *Principles of Biostatistics* (3rd Edition). Chapman and Hall/CRC. <https://doi.org/10.1201/9780429340512>
- Pineda-Tenor, D., Laserna-Mendieta, E. J., Timón-Zapata, J., Rodelgo-Jiménez, L., Ramos-Corral, R., Recio-Montealegre, A., & Reus, M. G. S. (2013). Biological variation and reference change values of common clinical chemistry and haematologic laboratory analytes in the elderly population. *Clinical Chemistry and Laboratory Medicine*, 51(4), 851–862. <https://doi.org/10.1515/CCLM-2012-0701/XML>
- Rifai, Nader., Chiu, R. W. K. ., Young, Ian., Burnham, C.-A. D. ., & Wittwer, C. . (2019). *Tietz Textbook of Clinical Chemistry and Molecular Diagnostics. 6 th ed. St. Louis, MO: Elsevier 2018*. Elsevier.
- Rothman, K. J., Greenland, S., & Associate, T. L. L. (2014). Modern Epidemiology. In *The Hastings Center report* (3rd Edition, Vol. 44). <http://www.ncbi.nlm.nih.gov/pubmed/24644503>
- Sackett, D. L., Rosenberg, W. M. C., Gray, J. A. M., Haynes, R. B., & Richardson, W. S. (1996). Evidence based medicine: what it is and what it isn't. *BMJ*, 312(7023), 71–72. <https://doi.org/10.1136/BMJ.312.7023.71>
- Schulz, K. F., Altman, D. G., & Moher, D. (2010). CONSORT 2010 statement: updated guidelines for reporting parallel group randomised trials. *BMJ (Clinical Research Ed.)*, 340(7748), 698–702. <https://doi.org/10.1136/BMJ.C332>
- Stevens, S. S. (1946). On the theory of scales of measurement. *Science*, 103(2684), 677–680. <https://doi.org/10.1126/SCIENCE.103.2684.677>;WEBSITE:WEBSITE:AAAS-SITE;JOURNAL:JOURNAL:SCIENCE;WGROU:STRING:PUBLICATION
- Westgard, J. O., Barry, P. L., Ehrmeyer, S. S., Plaut, D., Quam, E. F., Statland, B. E., por, T., & Migliarino, G. A. (2010). *Prácticas Básicas de Control de la Calidad*

3 a Edición Capacitación en Control Estadístico de la Calidad para Laboratorios Clínicos Con contribuciones de (3ra Edición). Westgard QC.
<https://colbiossa.com.ar/wp-content/uploads/2018/08/Practicas-Basicas-de-Control-de-la-Calidad-James-Westgard-1.pdf>

Young, D. S. (2002). Biological Variation: From Principles to Practice. *Clinical Chemistry*, 48(2), 395–395. <https://doi.org/10.1093/CLINCHEM/48.2.395A>