

UNIVERSIDAD VERACRUZANA

**INSTITUTO DE INVESTIGACIONES EN INTELIGENCIA ARTIFICIAL
DOCTORADO EN INTELIGENCIA ARTIFICIAL**



**Aprendizaje Supervisado en el Modelado y
Simulación Basados en Agentes**

Tesis

Que para obtener el grado de:

Doctor en Inteligencia Artificial

Presenta:

Mtro. Alejandro Platas-López

Director:

Dr. Alejandro Guerra-Hernández

Co-Director:

Dr. Nicandro Cruz-Ramírez

Xalapa-Enríquez, Veracruz, México

2023

1. INTRODUCCIÓN	1
1.1. Modelación	2
1.2. Definición del problema	5
1.3. Problema de investigación	7
1.4. Motivación	9
1.5. Hipótesis	11
1.6. Objetivos	13
1.6.1. General	13
1.6.2. Específicos	13
1.7. Justificación	14
1.8. Contribuciones	15
1.9. Organización del documento	15
1.10. Publicaciones	16
2. FUNDAMENTOS DE ABM Y ML	18
2.1. Modelado y Simulación Basado en Agentes	19
2.1.1. Complejidad y emergencia	20
2.1.2. Componentes	21
2.1.2.1. Agentes	22
2.1.2.2. Ambiente	22
2.1.2.3. Interacción	22
2.1.3. Validación y análisis	23
2.1.4. Ciclo de desarrollo	25
2.1.5. Descripción y formulación	26
2.1.5.1. GASS	28
2.1.5.2. TRACE	29
2.1.5.3. MIASE	31
2.1.5.4. MMRR	32



2.1.5.5.	CHEERS	33
2.1.5.6.	STRESS	34
2.1.5.7.	RAT-RS	35
2.1.5.8.	ODD	37
2.1.6.	Críticas a los ABM	38
2.1.6.1.	Computacionalmente caro	41
2.1.6.2.	Diseño	42
2.1.6.3.	Calibración	43
2.1.6.4.	Interpretación	44
2.1.6.5.	Necesidad de programación	45
2.2.	Aprendizaje Automático	46
2.2.1.	Aprendizaje	46
2.2.2.	Tipos de aprendizaje automático	48
2.2.2.1.	Supervisado	48
2.2.2.2.	No supervisado	50
2.2.2.3.	Por refuerzo	52
3.	REVISIÓN DE LA LITERATURA	55
3.1.	Metodología	56
3.1.1.	Objetivo de la revisión	56
3.1.2.	Organización de la revisión: temas principales	57
3.1.3.	Tipo de revisión de la literatura	58
3.1.3.1.	Revisión sistemática	58
3.1.3.2.	Revisión narrativa tradicional	59
3.1.4.	Búsqueda de literatura	59
3.1.5.	Criterios de inclusión o exclusión	60
3.2.	Revisión de la literatura	62
3.2.1.	Análisis bibliométrico	62
3.2.1.1.	Análisis estadístico inicial y bibliométrico	63
3.2.1.2.	Análisis de estructura conceptual	66
3.2.2.	Revisión narrativa	73
3.2.2.1.	Aprendizaje <i>offline</i> a priori	76
3.2.2.2.	Aprendizaje <i>offline</i> a posteriori	77
3.2.2.3.	Aprendizaje <i>online</i>	80
3.3.	Discusión	81
3.3.1.	Análisis de tendencias	82
3.3.2.	Oportunidades emergentes	83
3.3.3.	Estándares de documentación	85
3.4.	Conclusión	86
3.5.	Apéndice	89
4.	APRENDIZAJE OFFLINE EN MODELACIÓN BASADA EN AGENTES	112
4.1.	Propuesta de documentación	113
4.1.1.	Aprendizaje fuera de línea	114



4.1.2. A priori	117
4.1.3. A posteriori	117
4.2. Caso de estudio: Evasión fiscal	118
4.2.1. Aprendizaje supervisado a priori	122
4.2.2. Aprendizaje supervisado a posteriori	124
4.3. Apéndice	132
5. CONCLUSIÓN	176
ANEXO	181
A. Aprendizaje fuera de línea	182
A.1. El modelo de evasión de impuestos sobre la nómina	182
A.1.1. Visión general	183
A.1.2. Conceptos de diseño	184
A.1.3. Detalles	186
A.2. Resultados	193
A.2.1. Validación	194
A.2.2. El efecto del aprendizaje automático en la simulación	198
A.3. Conclusiones	199
BIBLIOGRAFÍA	220

ÍNDICE DE FIGURAS

1.1. Incremento en el número de publicaciones sobre ABM	10
1.2. Incremento en el número de publicaciones sobre ABM	11
1.3. Proporción de los artículos con títulos sobre ABM por área temática.	12
2.1. Ciclo de desarrollo de un modelo basado en agentes.	27
2.2. Estructura de descripciones de modelos siguiendo el protocolo ODD.	38
2.3. Críticas hacia los ABM y su correspondiente etapa de modelado	40
2.4. Representación del escenario general del aprendizaje por refuerzo.	53
3.1. Fórmula Scopus.	60
3.2. Fórmula Web of Science Core Collection.	60
3.3. Flujo de revisión de la literatura.	61
3.4. Tendencia de publicación en el área de ABM y ML	63
3.5. Red de citación conjunta entre revistas.	65
3.6. Principales autores de ABMS+ML	66
3.7. Documentos sobre ABMS+ML publicados por país.	67
3.8. Diagrama circular de flujo de colaboración entre países	68
3.9. Mapa temático.	69
3.10. Palabras clave más recurrentes	71
3.11. Evolución temática.	73
3.12. Artículos revisados sobre ML en ABM por etapa y algoritmo	74
3.13. Análisis del número de propuestas revisadas	82
3.14. Estructura del protocolo ODD	86
4.1. Estructura del ODD y elementos a actualizar	113
4.2. Principales Problemas Percibidos en México.	119
4.3. Diagrama de integración de ML en el Modelo de Evasión	123
4.4. Estructura de la red bayesiana aprendida con TAN	127
4.5. Resultados del proceso de optimización con el algoritmo JADE con Archivo. . .	130



A.1. Diagrama de integración de ML en el Modelo de Evasión de Impuestos sobre Nómina.	190
A.2. Interfaz gráfica de usuario del modelo.. . . .	194
A.3. La distribución de la nómina sigue una ley de potencias.	195
A.4. Skewness de la distribución de la nómina	195
A.5. Recaudación de impuestos por fuente de datos	196
A.6. RMSE de la recaudación de impuestos sobre la nómina predicha y real	197
A.7. Extensión de la evasión, impuestos recaudados y evasores en el sistema . . .	198

1.1. Principales críticas a los ABM	7
3.1. Principales fuentes editoriales de ABMS+ML	64
3.2. Documentos publicados por país del autor correspondiente	65
3.3. Artículos sobre ML en ABM revisados por etapa de modelado.	75
3.4. Publicaciones sobre ML para diseño de ABM	77
3.5. Publicaciones sobre ML para análisis de ABM	78
A.1. Variables de estado por tipo de agente	183
A.2. Variables de estado y método de inicialización	187
A.3. Input parameter initialization of baseline model.	188
A.4. Inicialización del parámetro de entrada del modelo de línea base.	188
A.5. Correlación de Spearman entre diferentes fuentes de datos	197

CAPÍTULO 1

INTRODUCCIÓN

Contenido del capítulo

1.1. Modelación	2
1.2. Definición del problema	5
1.3. Problema de investigación	7
1.4. Motivación	9
1.5. Hipótesis	11
1.6. Objetivos	13
1.6.1. General	13
1.6.2. Específicos	13
1.7. Justificación	14
1.8. Contribuciones	15
1.9. Organización del documento	15
1.10. Publicaciones	16

El modelado basado en agentes ha sido una herramienta fundamental en la investigación desde la década de los 70, iniciando con el influyente trabajo de [Schelling \(1971\)](#). Su popularidad y aplicación ha crecido significativamente desde la década de los 90, evidenciado por la aparición de congresos especializados en Europa ([Conte y Paolucci, 2014](#)). Su versatilidad ha permitido su uso en una amplia gama de áreas, desde las ciencias naturales hasta las ciencias sociales. Esta diversidad de aplicaciones es en gran medida gracias a sus ventajas comparativas frente a métodos alternativos, como aquellos basados en ecuaciones. Estas ventajas, que incluyen la heterogeneidad, autonomía, espacio explícito, interacciones locales, racionalidad acotada y dinámica de desequilibrio, se detallan en el capítulo 2.

Sin embargo, la literatura ha señalado algunos desafíos en su uso. Estos desafíos abarcan desde la demanda intensiva de recursos computacionales y la complejidad en su diseño, hasta la calibración de parámetros y la interpretación de resultados. Además, un reto adicional es la necesidad de habilidades de programación, las cuales pueden ser ajenas para investigadores no versados en ciencias computacionales.

La esencia de esta investigación es abordar y superar limitaciones en diseño, calibración y análisis de estos modelos. Para ello, se propone la integración de técnicas de aprendizaje automático (ML, por sus siglas en inglés) que enriquezcan etapas clave en el proceso de modelado y simulación basado en agentes, incluyendo el diseño, calibración, validación y análisis.

En este capítulo introductorio, se delinearán con profundidad los aspectos fundamentales del modelado y su relevancia en el contexto científico, las ventajas y desafíos de los modelos basados en agentes, la hipótesis central, los objetivos del estudio, y finalmente, la motivación, justificación y contribuciones clave de esta investigación.

1.1. Modelación

Entender la realidad no es una tarea sencilla ([Fischer, 1991](#)), para lidiar con ello, la ciencia recurre a teorías, datos y modelos para comprender e interpretar esta realidad ([Giere, 1999](#)), sin embargo, de acuerdo con [Giere](#) el razonamiento científico está en mayor medida basado en los modelos. Los modelos son una representación abstracta de esa realidad, y dependerá del propósito del científico, los elementos que deberán ser incluidos dentro del modelo o dejados fuera ([Johnson, 2001](#)).

En esta concepción, los modelos deben ser evaluados como herramientas diseñadas para propósitos específicos, es decir, la bondad de un modelo debe ser juzgada en función de lo bien que logra su propósito declarado ([Edmonds et al., 2019](#)). En este sentido, desde las ciencias sociales, [Epstein \(2008\)](#) identifica 17 diferentes objetivos de modelado, no obstante [Edmonds et al. \(2019\)](#) los resumen en las siguientes principales razones y sus respectivas definiciones:

- Predicción. Capacidad de anticipar de manera confiable aspectos bien definidos de

los datos que actualmente no se conocen con un grado útil de precisión a través de cálculos usando el modelo.

- Explicación. Establecer una posible cadena causal desde la configuración hasta sus consecuencias en términos de los mecanismos de una simulación.
- Descripción. Intento de representar parcialmente lo que es importante de un caso observado específico (o un pequeño conjunto de casos estrechamente relacionados).
- Exploración teórica. Establecer y luego caracterizar (o evaluar) hipótesis sobre el comportamiento general de un conjunto de mecanismos (utilizando una simulación).
- Ilustración. Una ilustración (usando una simulación) es para comunicar o aclarar una idea, teoría o explicación.
- Analogía. Una analogía (usando una simulación) es cuando los procesos ilustrados por una simulación se usan como una forma de pensar acerca de algo de manera informal.

Para alcanzar estos propósitos, existen diferentes enfoques de modelado empleados en las ciencias sociales. [Gilbert \(2020\)](#) identifica principalmente tres, a decir, el enfoque basado en variables usando ecuaciones estructurales; el enfoque basado en sistemas usando ecuaciones diferenciales y el enfoque basado en agentes.

Sin embargo, de acuerdo con el mismo [Gilbert](#), este último enfoque se ha vuelto popular ya que permite desarrollar modelos en donde las entidades individuales y sus interacciones son directamente representadas, además permite modelar heterogeneidad individual, representar explícitamente las reglas de comportamiento de los agentes y situar a los agentes en un espacio geográfico o de otro tipo. Permite a los modeladores representar múltiples escalas de análisis de una forma natural, la emergencia de estructuras a nivel macro o social desde las acciones individuales, y varios tipos de adaptación y aprendizaje. Ninguna de estas características es fácil de realizar con los otros enfoques de modelado.

En este mismo sentido, [Uhrmacher y Weyns \(2009\)](#) argumentan que la aplicación del enfoque de SMA ha sido reconocida como una metáfora útil para el modelado y simulación de sistemas complejos en un gran número de contextos, que van desde sistemas naturales a sociales, los cuales se caracterizan por la presencia de entidades autónomas cuyas la acción y la interacción en sus entornos físicos determina la evolución del sistema. A pesar

de ser un enfoque relativamente joven, comparado por ejemplo con el enfoque analítico basado en ecuaciones, se considera una de las perspectivas más exitosas para el modelado y simulación.

Aunque de acuerdo con [Dignum et al. \(2016\)](#), un modelo basado en agentes (del inglés, Agent-Based Model, ABM) es por definición, un modelo de un sistema multi-agente (del inglés, Multi-Agent System, MAS), históricamente las comunidades de investigación ABM y MAS han procedido sobre líneas independientes cercanas. Con el propósito de evitar los desafíos conceptuales y técnicos que supone la combinación de ambas perspectivas, la presente investigación será guiada principalmente por la línea de investigación de los modelos basados en agentes.

Como consecuencia del gran número de contextos en los que los ABM han sido empleados, se tienen identificados diferentes nombres para el mismo tipo de modelos, como por ejemplo simulación social ([Gilbert, 1997](#); [Conte y Terna, 2000](#)), sociedades artificiales ([Epstein y Axtell, 1997](#); [Doran, 1998](#); [Davidsson, 2001](#)), modelado basado en individuos ([Judson, 1994](#); [Grimm, 1999](#)), economía computacional basada en agentes ([Tesfatsion, 2002](#)), demografía computacional basada en agentes ([Billari et al., 2003](#)) o sociología (ciencia social) computacional basada en agentes ([Dessalles et al., 2008](#); [Squazzoni, 2012](#); [Anzola, 2019](#)). Para evitar confusión, todo este tipo de modelos serán tratados bajo el nombre general de “modelos basados en agentes” a lo largo de esta investigación.

Como se comentó previamente, se trata de un enfoque joven que data de la década de los 70, con el crecimiento del poder de cómputo y la accesibilidad a los ordenadores personales. Sin embargo, algunos investigadores ([Heath y Hill, 2008](#); [Gavin, 2018](#); [Klein et al., 2018](#); [Retzlaff et al., 2021](#)) concuerdan en situar la motivación de los ABM siglos más atrás, con la mano invisible de Adam Smith, la cual se establece que agentes individuales eligen acciones egoístas que derivan en ventaja mutua y beneficios sociales no intencionales para la comunidad; fenómenos sociales que son de hecho, características importantes que pueden ser representadas en los ABM.

Más allá del debate sobre el origen histórico conceptual de los modelos basados en agentes, la mayoría de los autores coinciden en las propiedades y ventajas que este enfoque de modelado ofrece en comparación con aproximaciones como el enfoque basado en

ecuaciones. Estas ventajas son resumidas por [Caiani y Caverzasi \(2017\)](#) de la siguiente manera:

- Heterogeneidad. Los agentes se modelan explícitamente y pueden diferir unas de otras por su dotación, preferencias, red local y cualquier característica que el modelador considere necesaria.
- Autonomía. No hay un controlador central y las regularidades emergen de manera autónoma de la interacción de los agentes.
- Espacio explícito. El espacio en el que las unidades actúan e interactúan (tanto unas con otras como con el sistema) se modela explícitamente. Esto nos permite definir los conceptos de “local” y “vecindario” (como en el ejemplo del modelo de segregación de [Schelling \(1969\)](#)).
- Interacciones locales. Las acciones de los agentes son locales ya que no interactúan con la totalidad del sistema, sino solo con sus vecinos, lo que puede estar referido a su posición espacial, económica o social.
- Racionalidad acotada. Los agentes basan sus decisiones en la heurística, ya que no están dotados ni de una información perfecta del funcionamiento del sistema en el que viven ni de una capacidad ilimitada para procesar la información disponible.
- Dinámica de desequilibrio. El equilibrio no se trata como un estado natural del sistema. La dinámica del desequilibrio es una preocupación central; en muchos casos, la dinámica del modelo no es ergódica debido a la relevancia de las condiciones iniciales y la dependencia de la trayectoria. Esto nos permite centrarnos en el proceso más que en el estado final y, por lo tanto, analizar sistemas en los que es posible que ni siquiera existan equilibrios.

1.2. Definición del problema

A pesar de las ventajas mencionadas, los modelos y simulaciones basadas en agentes no están exentos de críticas. Las principales críticas encontradas en la literatura se pueden agrupar en cinco rubros, a decir, el costo computacional, problemas para diseñar el modelo, inconvenientes en la calibración de parámetros, dificultades para interpretar la salida del modelo y finalmente, los conocimientos necesarios para programar el modelo. Estas críticas



se describirán en detalle en el capítulo 2 y son analizadas con el propósito de identificar áreas de mejora en la modelación basada en agentes.

En este sentido, se entiende por computacionalmente caro, la necesidad de ejecutar más de una corrida para generar conclusiones que sean estadísticamente significativas. Además, si se incrementa el número de agentes en el modelo el tiempo de ejecución también se incrementa, y en consecuencia los requerimientos de memoria también suben.

En cuanto a los problemas de diseño, nos referimos principalmente a dos, por un lado que los modeladores tienen gran flexibilidad en el diseño de sus modelos, que se ha caído en hacer modelos ad-hoc, con alto grado de complejidad, sin procurar buscar el nivel de detalle adecuado al fenómeno que se intenta reflejar. Por el otro lado, existe una preocupación sobre si los modelos construidos con correctos, dado que en ocasiones no se comprenden correctamente los mecanismos de comportamiento.

Por problema en la calibración de parámetros, podemos entender que algunos de los parámetros podrían ser sensibles a pequeñas variaciones, es decir, un mínimo cambio al inicio podría producir una salida totalmente diferente. Esto se vuelve un problema si se requiere una definición precisa o cuando se quiere entender como las entradas influyen en la salida del modelo. Sin embargo, no se tienen indicadores claros que describan la dinámica de incertidumbre del proceso entrada-salida. Esto se puede agravar si el nivel de detalle elegido en el modelo es erróneo, ya que puede llevar a modelos complejos y con muchos parámetros.

En lo que respecta al problema para interpretar o analizar el modelo, este se genera por que los patrones y los resultados de las interacciones de los agentes son intrínsecamente impredecibles, y no es fácil determinar que macro-proceso es la consecuencia de un determinado micro-proceso a nivel de agentes. Estos problemas permean sobre la visualización, problemas estadísticos relacionados con la definición de la cantidad de corridas y pruebas de hipótesis, la exploración espacial de soluciones y análisis de sensibilidad, así como la comunicación efectiva a audiencias no técnicas.

Finalmente, en lo que respecta al problema referido sobre la necesidad de saber programar, tiene que ver claramente con los requerimientos técnicos para implementar un modelo en algún lenguaje de programación. Pero a su vez se relaciona con el nivel de complejidad que

se desea en el modelo, está claro que un modelo sencillo será menos difícil de programar que uno con muchos procesos y tipos de agentes. En este sentido, la gran complejidad del comportamiento humano y el entorno urbano son extremadamente difíciles (o incluso imposibles) de capturar en un modelo informático, por lo que modelos informáticos nunca podrán dar cuenta de todo lo que puede afectar al sistema real. Por ello se deben tener claros los elementos relevantes del fenómeno a modelar.

De acuerdo con [Manzo \(2014\)](#), estas críticas se dan sobre todo cuando se ve en detalle la forma en que trabajan y se comparan con otros métodos que se auto-proclaman más fiables, como lo son los experimentos de laboratorio, los modelos matemáticos o los análisis estadísticos/econométricos. Sin embargo, en la apreciación de este autor, los detractores de la simulación basada en agentes (y los de la simulación en general) tienden a subestimar la presencia de problemas similares para los métodos que prefieren.

1.3. Problema de investigación

Tabla 1.1: Principales críticas a los modelos basados en agentes.

Autores	Críticas				
	Computacio- nalmente caro	Complejidad en diseño	Difícil calibración	Difícil interpretación	Problemas de programación
Axtell	✓				
Charteris et al.	✓	✓			
Breckling	✓	✓	✓		
Tesfatsion					✓
Li et al.				✓	
Nourqolipour y Shariff	✓				
Ligmann y Sun		✓	✓		
Malleson et al.		✓			✓
Rand y Rust	✓				
Miller et al.	✓				
Lengnick		✓		✓	
Conte y Paolucci	✓	✓			
Ghorbani et al.		✓		✓	
Kostadinov et al.	✓				
Siegfried	✓	✓	✓	✓	
Baldwin et al.				✓	✓
Lee et al.				✓	
Wilensky y Rand	✓				
Tarvid	✓	✓		✓	
Broniec et al.	✓				

Fuente: Elaboración propia.

La Tabla 1.1 agrupa las principales críticas encontradas. Los problemas que se han identificado sobre los modelos basados en agentes, se encuentran de alguna forma relacionados. Si partimos con las críticas respecto al diseño, pudimos observar que aunque para algunos investigadores, la flexibilidad que se tiene para implementar un modelo es una ventaja, esto también nos puede llevar a modelos más complejos, siempre y cuando la capacidad de programación del investigador lo permita, al ser más complejos, la cantidad de parámetros que se tienen incrementa, aumentando en consecuencia el costo computacional, por supuesto, esto además puede conducir a modelos más complejos de analizar y validar.

Lo anterior se relaciona con lo mencionado por Siegfried (2014), quien menciona que una mala decisión en el nivel de detalle en el modelo, puede desencadenar en modelos con alto consumo computacional y posiblemente con calibraciones erróneas. Y a su vez dificultará su análisis, porque será difícil distinguir qué comportamientos a nivel agente producirán un resultado a nivel sistema.

Aunque la capacidad de programación, es un límite natural importante para hacer crecer la complejidad de un modelo basado en agentes, e inclusive para analizar el modelo mismo, la capacidad de programación no será considerado como un problema de investigación que deba atenderse para mejorar el desarrollo e implementación de modelos basados en agentes, pero si será visto como una atenuante que puede incidir de alguna manera en cada etapa de modelado.

Por lo tanto, el resto de críticas relativas al diseño, costo computacional, calibración e interpretación, se consideran el problema de investigación que será atendido. Este subconjunto de críticas hacen referencia a alguna etapa específica del proceso de modelación, a decir, diseño, calibración, exploración-experimentación y/o análisis-validación.

Pregunta de investigación

¿Cómo puede el aprendizaje supervisado mejorar los procesos de diseño, calibración y análisis en el Modelado y Simulación Basados en Agentes, y cuál sería el marco de trabajo para integrarlos de manera efectiva?

1.4. Motivación

De acuerdo con [Edmonds \(2000\)](#) existe un deseo comprensible de analizar la complejidad de los sistemas naturales (o sociales) en lugar de solo los modelos de esos sistemas, pero si los sistemas naturales tienen niveles inherentes de complejidad, pueden estar más allá de nuestra capacidad de modelar, pero en muchas ocasiones también van más allá de nuestro interés científico. En la práctica, no existe un límite superior en su complejidad, ya que siempre se pueden considerar con más detalle o incluyendo más aspectos.

Por lo tanto, la modelación es una práctica fundamental en la ciencia ([Schwarz et al., 2009](#)). Un modelo científico es una representación abstracta y simplificada de un sistema de fenómenos que hace que sus características centrales sean explícitas y visibles y se puede utilizar para generar explicaciones y predicciones.

Para ello, los científicos de diferentes disciplinas han recurrido a diferentes alternativas de modelado. Uno de estos enfoques, expuesto anteriormente y relativamente nuevo, ha surgido como uno de los más promisorios y el interés generado en la comunidad científica ha crecido significativamente. Para evidenciarlo, se efectuó una consulta en la base de datos bibliográfica Scopus con la formula siguiente:

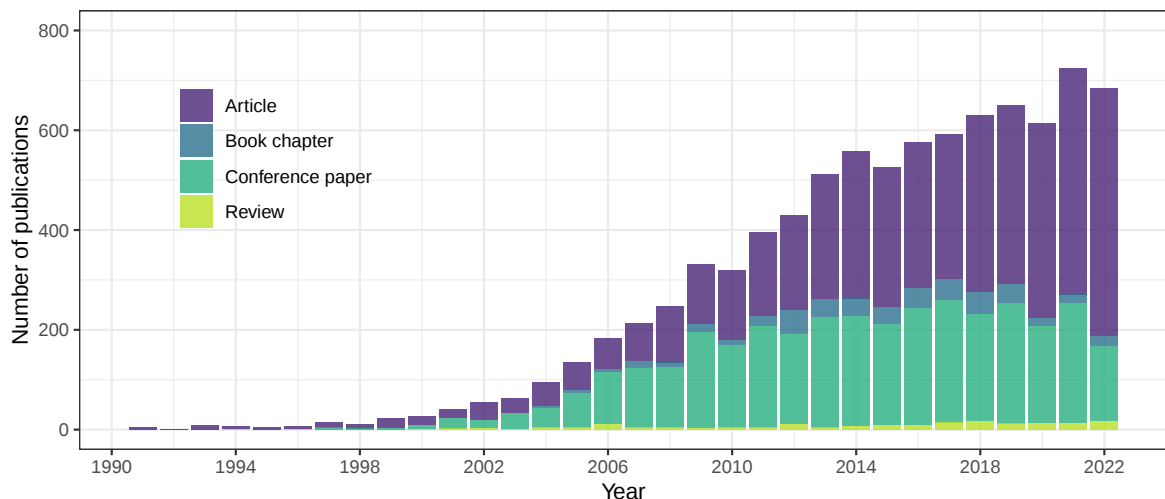
Fórmula Scopus

```
TITLE(((("individual-based") OR ("agent-based")) AND ((model* OR simulat*)  
OR (comput*)) OR ("artificial society") OR ("social simula*"))  
AND PUBYEAR >1979
```

La consulta en Scopus generó como resultado, un total de 11,538 artículos. Como se aprecia en la Figura 1.1, el número de publicaciones con títulos sobre modelado basado en agentes se ha incrementado con el paso de los años. De estas publicaciones, 5,779 (50.1 %) son artículos de revista, 4,866(42.2 %) son artículos de conferencia, 477 (4.1 %) capítulos de libro y 222 (1.9 %) revisiones de literatura.

Originalmente, alrededor de la década de los 90's, este tipo de modelos comenzó a popularizarse en el área de ecología con modelos de población de animales ([Grimm, 1999](#)), en lo que denominaron modelos basados en individuos. Sin embargo, a partir de la década de los 2000, estos modelos comenzaron a extenderse en otras disciplinas bajo el nombre

Figura. 1.1: Incremento en el número de publicaciones indexadas a Scopus sobre modelación basada en agentes y modelación basada en individuos.



Fuente: Platas-López et al. (2023b)

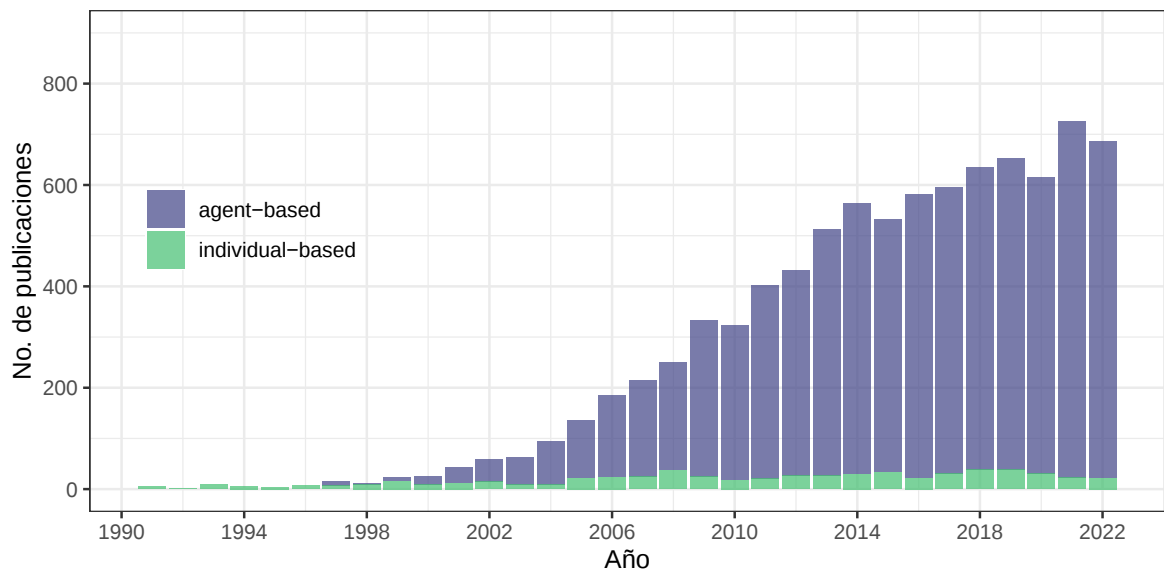
de modelos basados en agentes. La Figura 1.2 evidencia estos progresos a lo largo de las últimas tres décadas.

Es difícil definir con precisión en cual de las diferentes disciplinas científicas se han efectuado más aportaciones, ya que en algunos artículos intervienen dos o más áreas en mayor o menor medida. En la Figura 1.3 se realiza una aproximación a la clasificación de publicaciones por áreas temáticas. De tal forma que se puede apreciar que en las que más se ha aplicado son las ciencias sociales (1,485), ciencias medioambientales (1,251), ciencias agrícolas y biológicas (653), negocios, gestión y contabilidad (626) y economía, econometría y finanzas (511).

Sin embargo, a pesar de su utilidad y ventajas en diversas áreas, los modelos basados en agentes no son una panacea y como se describió en la sección anterior, son susceptibles de mejora. La motivación principal de esta investigación radica en superar las siguientes limitaciones: problemas con el diseño, calibración, la validación y el análisis.

Para lograrlo, se ha planteado la incorporación de técnicas de aprendizaje automático que beneficien algunas de las etapas del proceso de implementación del modelado y simulación basados en agentes, ya sea la especificación, la simulación, la calibración e inclusive la validación.

Figura. 1.2: Incremento en el número de publicaciones indexadas a Scopus sobre modelación basada en agentes y modelación basada en individuos.



Fuente: Elaboración propia con datos de Scopus.

1.5. Hipótesis

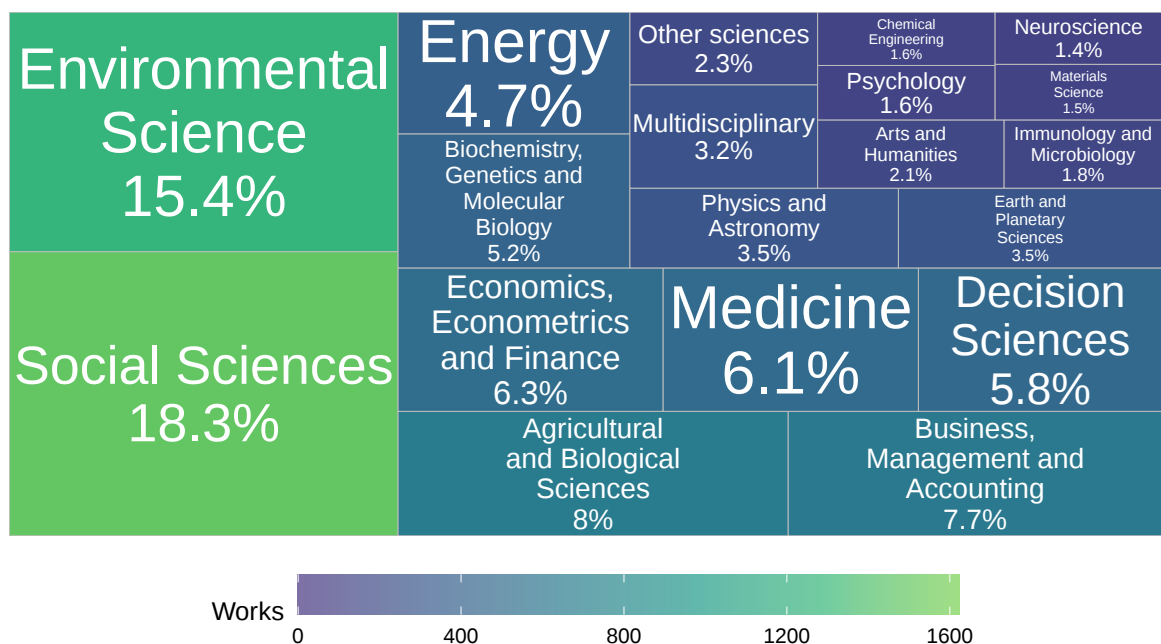
Hipótesis

La integración de técnicas de aprendizaje supervisado en la modelación y simulación basadas en agentes puede abordar y mejorar desafíos claves asociados con el diseño, calibración y análisis de estos modelos. Además, un flujo de trabajo estructurado y adaptado a partir del protocolo ODD puede proporcionar una guía clara para aquellos investigadores y profesionales que buscan combinar aprendizaje automático supervisado con ABM.

Para validar esta hipótesis, se emprendieron una serie de pasos metodológicos. Primero, se realizó una revisión exhaustiva de la literatura, identificando los problemas recurrentes y las soluciones propuestas al incorporar aprendizaje automático, en particular aprendizaje supervisado, en la modelación basada en agentes. Esta revisión permitió corroborar que ya existen esfuerzos previos en esta dirección, aunque con enfoques y flujos de trabajo variados.

Luego, se desarrolló un caso de estudio en el que se integraron técnicas de aprendizaje supervisado al proceso de simulación. Esto no solo demostró la factibilidad de la integración,

Figura. 1.3: Proporción de los artículos con títulos sobre ABM por área temática.



Fuente: Platas-López et al. (2023b)

sino también el potencial del aprendizaje supervisado para mejorar aspectos clave del diseño del modelo basado en agentes. En particular, se destacó la eficiencia de la utilización de datos existentes para aprender aspectos del fenómeno a ser representado.

En cuanto al análisis y la interpretación de los resultados, se aplicaron técnicas de aprendizaje automático de caja blanca, como las redes bayesianas, para ofrecer transparencia en el entendimiento de los modelos. Esto demostró ser especialmente útil, ya que ofreció una visión más clara y detallada de los patrones y comportamientos emergentes dentro de la simulación.

Finalmente, con base en la revisión literaria y la experiencia adquirida en el desarrollo del caso de estudio, se propuso un flujo de trabajo adaptado a partir del protocolo ODD. Este nuevo marco busca ser una herramienta valiosa para guiar a otros investigadores y profesionales en la integración de técnicas de aprendizaje supervisado en la modelación basada en agentes.

En resumen, a través de estos esfuerzos y el desarrollo práctico, se pudo constatar que la hipótesis inicial tiene fundamento. Las técnicas de aprendizaje supervisado no solo son

compatibles con el modelado basado en agentes, sino que tienen el potencial de mejorar significativamente el diseño, calibración y análisis de estos modelos.

1.6. Objetivos

La aplicación de aprendizaje supervisado en Modelos Basados en Agentes presenta un panorama prometedor para abordar las críticas y desafíos tradicionales asociados con estos modelos. El propósito principal de esta investigación radica en la exploración, evaluación y propuesta de métodos que puedan hacer esta integración más efectiva y accesible para la comunidad científica y de desarrollo. Al mismo tiempo, se busca establecer un marco de referencia que no solo documente este proceso integrativo sino que también lo facilite y promueva. Para lograr una comprensión integral del escenario actual y las posibilidades futuras, esta tesis se propone objetivos bien definidos, tanto generales como específicos, que guiarán el desarrollo del estudio y permitirán una implementación práctica y analítica del aprendizaje supervisado en ABM.

1.6.1. General

Evaluar y proponer métodos basados en aprendizaje supervisado para mejorar el diseño, calibración y análisis de Modelos Basados en Agentes, estableciendo un flujo de trabajo efectivo que pueda ser adoptado en la comunidad de investigación y desarrollo.

1.6.2. Específicos

O.E. 1. Revisar críticamente la literatura existente para identificar y categorizar esfuerzos anteriores que han intentado integrar el aprendizaje automático con el modelado y simulación basados en agentes.

O.E. 2. Proponer un marco de referencia que facilite la documentación y entendimiento de la integración de técnicas de aprendizaje supervisado en el modelado y simulación basados en agentes.

O.E. 3. Implementar un caso de estudio que encapsule un modelo de aprendizaje supervisado dentro de una simulación, evaluando su eficacia y aplicabilidad.

O.E. 4. Aplicar técnicas de caja blanca, para analizar y entender de forma transparente los resultados del modelo.

1.7. Justificación

La adopción de técnicas de aprendizaje supervisado en Modelos Basados en Agentes surge como una respuesta innovadora a los desafíos tradicionales asociados con el diseño, calibración y análisis de estos modelos. Esta tesis nace de la observación de que, a pesar de los esfuerzos aislados en la literatura para resolver desafíos específicos de ABM con aprendizaje automático, aún existe una brecha en la consolidación de un flujo de trabajo estandarizado y optimizado.

La justificación de este estudio reside en dos pilares principales:

Explosión de Datos: La era contemporánea ha sido testigo de una proliferación masiva de conjuntos de datos detallados y ricos. Estos conjuntos de datos, que describen una multitud de fenómenos con precisión sin precedentes, son un caldo de cultivo para técnicas de aprendizaje automático. A través del aprendizaje supervisado, es posible destilar patrones complejos de estos datos y aplicarlos en el diseño de ABM, abordando de este modo la complejidad inherente en la creación de estos modelos. Además, la producción de datos generados por simulaciones puede beneficiarse enormemente de técnicas de aprendizaje automático para su calibración, validación y análisis posterior, optimizando así el proceso de interpretación y garantizando resultados más fiables.

Naturaleza Inherente de ABM: No siempre se dispone de un entendimiento total sobre el fenómeno en estudio al desarrollar un ABM. En tales casos, el aprendizaje automático, especialmente el enfoque supervisado, puede actuar como un puente, complementando la simulación con datos y reglas derivadas. Aunque la tesis no se adentra en el aprendizaje en línea, su mención destaca la capacidad del aprendizaje automático para aprender activamente de los estados emergentes en la simulación, lo que enriquece aún más el modelo y facilita un análisis robusto.

Por último, al tomar como referencia el protocolo ODD, esta tesis busca adaptarlo y expandirlo para abarcar la creciente intersección de aprendizaje automático y ABM, abogando por un marco más inclusivo que guíe a los investigadores en este territorio interdisciplinario.

1.8. Contribuciones

La presente investigación arroja luz sobre la intersección del aprendizaje supervisado y el modelado basado en agentes. Las contribuciones principales de este trabajo doctoral son:

- **Revisión Crítica de ABM:** Se presentaron argumentos y críticas contra la naturaleza de ABM, en particular, los desafíos relacionados con el diseño, calibración, validación e interpretación.
- **Clasificación de Esfuerzos ML-ABM:** A través de una revisión detallada de la literatura, este trabajo categoriza y ordena los esfuerzos previos en la aplicación de aprendizaje automático para resolver problemas específicos en ABM en dos clases principales: aprendizaje offline y aprendizaje online.
- **Protocolo ODD Modificado:** Proporcionamos una adaptación del protocolo ODD, uno de los estándares en ABM, para acomodar especificaciones y flujos de trabajo de ABM que integran aprendizaje automático supervisado, abordando la brecha actual en ODD que parece centrarse más en el aprendizaje online.
- **Aplicación Práctica a través de Casos de Estudio:** La teoría presentada en esta tesis se ejemplifica y valida a través de casos de estudio significativos. Estos casos ilustran tanto la encapsulación de un modelo de aprendizaje automático dentro de una simulación (aprendizaje a priori) como la aplicación de técnicas de aprendizaje automático supervisado para el análisis posterior a la simulación (aprendizaje a posteriori).
- **Innovaciones en Técnicas de Aprendizaje:** En particular, se introduce una metodología de aprendizaje discriminativo para optimizar la Conditional Log Likelihood en la construcción de redes bayesianas, aportando una precisión mejorada en la clasificación sin sacrificar la transparencia.

1.9. Organización del documento

La investigación se encuentra organizada de la siguiente manera. El Capítulo 2 aborda dos áreas clave de la investigación: el Modelado y Simulación Basado en Agentes y el Aprendizaje Automático. Se detallará la estructura, evolución, y desafíos del ABM, incluyendo el Protocolo ODD, mientras que en ML se definirán sus tipos y enfoques principales, centrando la atención

en el aprendizaje supervisado. La finalidad es sentar una base teórica para la integración futura de ambas disciplinas.

En el Capítulo 3 se ofrece una revisión profunda del crecimiento interdisciplinario entre el aprendizaje automático y la modelación basada en agentes, en la cual se han realizado numerosas investigaciones. Se destacan tendencias, lagunas y oportunidades futuras. Iniciando con la metodología de búsqueda, se detalla una revisión exhaustiva que combina enfoques narrativos y sistemáticos, culminando con un análisis del estado actual del campo y potenciales direcciones para futuros estudios.

El Capítulo 4 aborda el aprendizaje offline en ABM, definiendo su marco dentro del protocolo ODD y explorando su utilidad tanto previo a la simulación como en el análisis de resultados. Se presenta un ejemplo práctico de cómo el aprendizaje supervisado puede ser integrado en la simulación, facilitando el diseño y ofreciendo análisis profundos de los resultados. Finalmente en el Capítulo 5 se dan algunas reflexiones finales y trabajo futuro.

1.10. Publicaciones

En esta sección, se enlistan las publicaciones de revista derivadas de la presente investigación. Cada entrada especifica la revista de publicación, con una notación referente al Factor de Impacto de acuerdo con el Journal Citation Reports, una métrica que cuantifica la visibilidad y el impacto de las revistas académicas. Asimismo, se señala la relación directa con el capítulo correspondiente en la tesis doctoral para facilitar su referencia y contextualización. Estas publicaciones reflejan la profundidad y el alcance de la investigación realizada y sirven como una base sólida para futuras investigaciones en el campo.

1. Alejandro Platas-López, Efrén Mezura-Montes, Nicandro Cruz-Ramírez y Alejandro Guerra-Hernández, **Discriminative learning of bayesian network parameters by differential evolution**, *Applied Mathematical Modelling*, 2021, <https://doi.org/10.1016/j.apm.2020.12.026>. **Q1, IF 5.0**. Relacionado con el Capítulo 4.
2. Alejandro Platas-López, Alejandro Guerra-Hernández, Francisco Grimaldo, Nicandro Cruz-Ramírez, Efrén Mezura-Montes y Marcela Quiroz-Castellanos. **dplbnDE: An R package for discriminative parameter learning of Bayesian Networks by Diffe-**



- rential Evolution**, *SoftwareX*, 2023, <https://doi.org/10.1016/j.softx.2023.101442>. **Q2, IF 3.4**. Relacionado con el Capítulo 4.
3. Alejandro Platas-López, Alejandro Guerra-Hernández, Marcela Quiroz-Castellanos y Nicandro Cruz-Ramírez, **A survey on agent-based modelling assisted by machine learning**, *Expert Systems*, 2023, <https://doi.org/10.1111/exsy.13325>. **Q2, IF 3.3**. Relacionado con el Capítulo 3
 4. Alejandro Platas-López, Alejandro Guerra-Hernández, Marcela Quiroz-Castellanos y Nicandro Cruz-Ramírez, **Agent-Based Models Assisted by Supervised Learning: A Proposal for Model Specification**, *Electronics*, 2023, <https://doi.org/10.3390/electronics12030495>. **Q2, IF 2.9**. Relacionado con el Capítulo 4.

CAPÍTULO 2

FUNDAMENTOS DE ABM Y ML

Contenido del capítulo

2.1. Modelado y Simulación Basado en Agentes	19
2.1.1. Complejidad y emergencia	20
2.1.2. Componentes	21
2.1.2.1. Agentes	22
2.1.2.2. Ambiente	22
2.1.2.3. Interacción	22
2.1.3. Validación y análisis	23
2.1.4. Ciclo de desarrollo	25
2.1.5. Descripción y formulación	26
2.1.5.1. GASS	28
2.1.5.2. TRACE	29
2.1.5.3. MIASE	31
2.1.5.4. MMRR	32
2.1.5.5. CHEERS	33
2.1.5.6. STRESS	34
2.1.5.7. RAT-RS	35
2.1.5.8. ODD	37
2.1.6. Críticas a los ABM	38
2.1.6.1. Computacionalmente caro	41
2.1.6.2. Diseño	42
2.1.6.3. Calibración	43
2.1.6.4. Interpretación	44
2.1.6.5. Necesidad de programación	45
2.2. Aprendizaje Automático	46
2.2.1. Aprendizaje	46
2.2.2. Tipos de aprendizaje automático	48
2.2.2.1. Supervisado	48
2.2.2.2. No supervisado	50
2.2.2.3. Por refuerzo	52

El Modelado y Simulación Basado en Agentes y el Aprendizaje Automático representan dos áreas prominentes en la investigación computacional y la modelización de sistemas. Ambas técnicas, si bien tienen distintos enfoques, ofrecen herramientas metodológicas y analíticas de gran relevancia para la comprensión y modelación de sistemas complejos.

En la primera sección de este capítulo, se examinará en detalle el Modelado Basado en Agentes. Se iniciará con una descripción del ciclo de desarrollo de los modelos basados en agentes, proporcionando una base para comprender la estructura y evolución de estas simulaciones. Posteriormente, se abordarán conceptos esenciales como la complejidad, emergencia y los componentes primordiales de los sistemas ABM: agentes y ambiente. A continuación, se presentará el Protocolo ODD, una metodología establecida para la descripción y formulación de estos modelos. Finalmente, se discutirán diversas críticas y limitaciones asociadas a la ABM, para ofrecer una perspectiva completa sobre sus desafíos y áreas de mejora.

La segunda sección se centra en el Aprendizaje Automático. Se proporcionará una definición formal de “aprendizaje” en el contexto de las máquinas y se expondrán los diferentes tipos de aprendizaje automático. Si bien el enfoque principal será el aprendizaje supervisado, también se presentarán brevemente enfoques alternativos como el aprendizaje no supervisado y el aprendizaje por refuerzo.

Este capítulo busca establecer una base teórica sólida sobre ABM y ML, preparando al lector para los capítulos subsiguientes donde se explorará en profundidad la interrelación entre estas dos áreas.

2.1. Modelado y Simulación Basado en Agentes

Un agente, en el contexto computacional, se define como una entidad autónoma con características y acciones definidas. El modelado basado en agentes se emplea para representar fenómenos mediante la interacción de dichos agentes en un ambiente específico. En agentes inteligentes, el aprendizaje puede ser concebido como un proceso de adaptación basado en retroalimentación para optimizar su desempeño.

Según [Railsback y Grimm \(2019\)](#), en la historia de la modelación científica, la complejidad de los modelos estuvo vinculada a su factibilidad matemática. Antes de la era computacional,

la modelación dependía principalmente del cálculo diferencial, lo que llevó a simplificar los modelos para poder resolverlos. Sin embargo, esta simplificación restringía la modelación a problemas menos complejos.

La llegada de la simulación computacional eliminó esas restricciones matemáticas. Se inició una transición hacia modelos menos restringidos y más representativos de sistemas reales. Los modelos basados en agentes se distinguen por su capacidad de representar componentes individuales y sus comportamientos. En vez de describir un sistema en su totalidad, se detalla la dinámica de sus agentes individuales.

En estos modelos, los agentes son concebidos como entidades autónomas y distintas que interactúan en su entorno local. Estos pueden representar desde organismos y humanos hasta empresas o instituciones con objetivos concretos. Los agentes son variables en características y generalmente interactúan solo con sus vecinos cercanos, ya sea en un espacio físico o una representación abstracta como una red.

Estos agentes adaptan su comportamiento en función de sus propios estados, el de otros agentes y las condiciones del ambiente. La modelación basada en agentes es especialmente relevante en sistemas donde la emergencia juega un papel crucial, es decir, donde las dinámicas del sistema se derivan de las interacciones individuales.

La principal ventaja de los modelos basados en agentes radica en su naturaleza multi-nivel. Mientras que algunos enfoques modelan solamente sistemas o agentes por separado, el ABM integra ambos niveles y sus interacciones, permitiendo un análisis bidireccional: comprender las acciones individuales en el contexto global y viceversa.

2.1.1. Complejidad y emergencia

Vivimos en una era de interconexiones y complejidades sin precedentes. En tiempos anteriores, la naturaleza limitada y localizada de la información y la tecnología permitía a los individuos y sociedades centrarse en sistemas más simples y relaciones lineales. Sin embargo, el progreso tecnológico y la globalización han cambiado radicalmente esta perspectiva, haciendo que los sistemas complejos sean no sólo más visibles, sino también más influyentes en nuestras vidas diarias.

Considere, por ejemplo, cómo un solo evento en una parte del mundo puede tener

efectos ondulantes que afectan a casi todos los aspectos de la vida en otra. Un brote de enfermedad en una ciudad puede, en cuestión de semanas o incluso días, afectar a la economía global, a las políticas de salud pública en países distantes y a las interacciones sociales en comunidades que antes parecían inmunes a tales influencias externas. Estos son ejemplos de emergencia, donde propiedades o fenómenos nuevos y a menudo inesperados surgen de interacciones simples entre agentes en un sistema.

Según [Wilensky y Rand \(2015\)](#), esta no es una idea nueva, pero ha ganado una atención sin precedentes en la era moderna. Los sistemas complejos, definidos como entidades compuestas por múltiples agentes interactuantes que operan de manera distribuida, se han convertido en el foco de una amplia variedad de disciplinas. Desde economistas que intentan predecir las fluctuaciones del mercado, hasta ecologistas que estudian ecosistemas dinámicos, la ciencia de los sistemas complejos ha ofrecido nuevas lentes para examinar y comprender nuestro mundo interconectado.

Este auge en el estudio de los sistemas complejos durante los años 90 se debió, en parte, a la convergencia de disciplinas que antes operaban en silos. La comprensión de que ciertos principios, como la adaptación, la emergencia y la auto-organización, eran aplicables en múltiples campos llevó a una colaboración más estrecha y al desarrollo de nuevas herramientas y métodos de análisis.

Dentro de esta amplia paleta de herramientas, la modelación basada en agentes ha emergido como una metodología particularmente poderosa. Los ABM permiten a los investigadores diseñar, experimentar y analizar sistemas en los que los agentes individuales, con sus propias reglas y comportamientos, interactúan entre sí y con su entorno. Estos modelos capturan la esencia de la complejidad y emergencia, ya que permiten que fenómenos complejos surjan de interacciones simples, replicando así muchos de los sistemas y fenómenos que observamos en el mundo real.

2.1.2. Componentes

Los sistemas complejos se pueden modelar y explicar productivamente creando agentes y ambientes, describiendo su comportamiento a través de reglas de agentes y especificando interacciones agente-agente y agente-entorno. Esta descripción es una imagen simplificada

de un ABM. Los componentes principales de cualquier ABM son los agentes, el ambiente y las interacciones. A continuación se describen a detalle cada uno de estos componentes.

2.1.2.1. Agentes

Los agentes son las unidades básicas del modelado basado en agentes. Como tal, es importante elegir cuidadosamente el diseño de sus agentes. Los dos aspectos principales que definen a los agentes son las propiedades que tienen y las acciones (a veces llamadas comportamientos o métodos) que pueden ejecutar. Las propiedades de los agentes son el estado interno y externo de los agentes: sus datos y su descripción. Los comportamientos o acciones de los agentes son lo que los agentes pueden hacer.

2.1.2.2. Ambiente

El entorno consiste en las condiciones y hábitats que rodean a los agentes a medida que actúan e interactúan dentro del modelo. El entorno puede afectar las decisiones de los agentes y, a su vez, puede verse afectado por las decisiones de los agentes. Existen dos tipos principales de ambientes, los espaciales y los basados en red.

Antes de hablar de los tipos de entornos, es importante mencionar que el propio entorno se puede implementar de varias formas. En primer lugar, el entorno puede estar compuesto de agentes de modo que cada parte individual del entorno pueda tener un conjunto completo de propiedades y acciones. Esto permite que diferentes partes del medio ambiente tengan diferentes propiedades y actúen de manera diferente en función de sus interacciones locales.

Un segundo enfoque representa el medio ambiente como un gran agente, con un conjunto global de propiedades y acciones. Otro enfoque más es implementar el entorno fuera del conjunto de herramientas del ABM. Por ejemplo, podría ser manejado por un conjunto de herramientas de sistemas de información geográfica (GIS), o podría ser manejado por un conjunto de herramientas de análisis de redes sociales, y el ABM puede interactuar con ese entorno.

2.1.2.3. Interacción

Las interacciones son el nexo fundamental entre agentes y entre agentes y su entorno. Sin interacción, un ABM no sería más que una serie de agentes individuales actuando en

aislamiento, lo cual no reflejaría con precisión la mayoría de los sistemas complejos que se pretenden modelar. Las interacciones dan lugar a patrones emergentes y comportamientos dinámicos, que son esenciales para entender la complejidad de estos sistemas.

Las interacciones pueden ser de varios tipos, tales como:

Agente-Agente Estas son las interacciones entre dos o más agentes. Pueden ser colaborativas, en las que los agentes trabajan juntos para lograr un objetivo común; competitivas, donde los agentes trabajan en oposición entre sí; o neutras, donde los agentes simplemente se afectan mutuamente sin intenciones claras de cooperar o competir.

Agente-Entorno Aquí, los agentes interactúan con su entorno circundante. Esta interacción puede ser pasiva, como un agente moviéndose a través de un espacio, o activa, donde el agente cambia o se ve influenciado por el entorno, como un animal buscando alimento.

Es fundamental para el diseño y el análisis de un ABM entender cómo y por qué los agentes interactúan. La estructura y naturaleza de estas interacciones pueden influir significativamente en los resultados del modelo. Por ejemplo, un modelo donde los agentes compiten ferozmente por recursos limitados se comportará muy diferente a uno donde los agentes colaboran y comparten recursos. Del mismo modo, la forma en que los agentes interactúan con su entorno puede determinar cómo evoluciona y cambia ese entorno con el tiempo.

2.1.3. Validación y análisis

También es importante enfatizar temas claves como la calibración, validación empírica, validación local, y el análisis de sensibilidad. Comencemos con la validación, que según [Wikstrom y Nelson \(2022\)](#) se refiere a evaluar si los componentes del modelo se comportan como se espera y si los resultados se corresponden con fenómenos observados.

Al profundizar en la validación, [Fagiolo et al. \(2019\)](#) distinguen dos tipos principales: la validación estructural, que evalúa cómo la simulación representa el modelo conceptual del sistema real, y la validación de salida, que evalúa si la salida de la simulación refleja los comportamientos históricos del sistema real. En este sentido, [Windrum et al. \(2007\)](#) critican la falta de consenso entre los economistas en la validación empírica de ABM y proponen desarrollar un protocolo mínimo común para el análisis de estos modelos, enfatizando la necesidad de abordar la heterogeneidad y las cuestiones metodológicas.

En relación a la vinculación del modelo con los datos, [Guerini y Moneta \(2017\)](#) exponen su relevancia, argumentando que la calibración y replicación de propiedades estadísticas son pasos iniciales, pero insuficientes. Coincidiendo con [Fagiolo et al. \(2019\)](#), afirman que la confiabilidad se puede mejorar mediante un ejercicio de validación diseñado para proporcionar evidencia de que el mecanismo generador de datos del modelo es una representación adecuada del mecanismo real.

En referencia a la validez de los modelos, [Tieleman \(2021\)](#) enfatiza que los modelos deben ser vistos como instrumentos contruidos para un propósito específico y deben validarse dentro del dominio del modelo. Esto es coherente con la clasificación de niveles de validez empírica que [Fagiolo et al. \(2019\)](#) referencian, donde los ABM se categorizan en niveles de 0 a 3 según su alineación con estructuras empíricas macro y micro.

En cuanto al análisis de sensibilidad, es una herramienta clave para discernir por qué los modelos producen los resultados que hacen y si otras suposiciones arrojaran resultados similares. [Borgonovo et al. \(2022\)](#) proponen un enfoque sistemático para llevar a cabo análisis de sensibilidad en ABM, especialmente teniendo en cuenta elementos no paramétricos. [Thiele et al. \(2014\)](#) también proporcionan herramientas y métodos para la calibración de parámetros y el análisis de sensibilidad en ABM sugiriendo el uso de R.

Un enfoque de análisis de sensibilidad basado en varianza es propuesto por [Saltelli et al. \(2010\)](#), el cual es efectivo para describir patrones de sensibilidad en modelos con factores de entrada inciertos. En este contexto, es importante mencionar los métodos de Coeficientes de Regresión Estandarizados [Saltelli et al. \(2007, p. 18\)](#) que son índices de sensibilidad basados en supuestos lineales o monótonos en el caso de factores independientes, los Coeficientes de Correlación Parciales ([Clouvel et al., 2023](#)) que son medidas basadas en la varianza basadas en suposiciones lineales, en el caso de factores correlacionados, y la Estimación de efectos de Shapley basada en datos a través de vecinos más cercanos ([Broto et al., 2020](#)) que implementa la estimación de varios índices de sensibilidad utilizando solo evaluaciones de modelos a través de la clasificación o búsqueda de vecinos más cercanos, lo que permite trabajar con diferentes tipos de entradas, entradas dependientes y múltiples salidas.

La validación y el análisis de sensibilidad son componentes críticos en la modelación basada en agentes. Para el caso específico de la evasión fiscal en México, es importante asegurar que los agentes actúen como deberían a nivel local y que los modelos represen-

ten adecuadamente las realidades socioeconómicas y estructurales del sistema tributario. Además, teniendo en cuenta la percepción de inseguridad y corrupción como factores institucionales, es necesario que la calibración y validación de modelos estén alineadas con datos empíricos y que se aplique un riguroso análisis de sensibilidad para evaluar la robustez de los resultados.

Finalmente, una consideración importante en el desarrollo de ABM es cómo aprovechar la existencia y generación de datos para asistir en el diseño y análisis del modelo respectivamente. El ML ha emergido como una herramienta valiosa en este aspecto. No solo permite incorporar datos reales en la fase de diseño para crear modelos más precisos, sino que también se puede utilizar en la fase de análisis para identificar patrones y relaciones complejas que podrían no ser evidentes a través de análisis convencionales. Dicho esto, el análisis de un ABM puede ser complicado debido a la naturaleza altamente dinámica y no lineal de estos modelos. El gran número de interacciones y la emergencia de comportamientos inesperados requieren herramientas y enfoques analíticos sofisticados.

2.1.4. Ciclo de desarrollo

A este punto se hace evidente la importancia de una metodología estructurada para llevar a cabo un modelado eficiente y científicamente sólido. Uno de los primeros pasos en esta metodología implica la conceptualización del modelo, donde las fases de diseño, experimentación, validación y análisis son cruciales.

Un enfoque iterativo en el modelado basado en agentes es necesario, tal y como destacan [Railsback y Grimm \(2019\)](#). Ellos sugieren un “ciclo de modelado” que comprende varias tareas, como formular la pregunta de investigación, recopilar hipótesis, elegir escalas y entidades, implementar el modelo, y finalmente, analizar, probar y revisar el modelo. Este enfoque cíclico enfatiza que el modelo inicial siempre puede mejorarse, y puede haber necesidad de visitar y refinar ciertos aspectos del modelo. Además, [Railsback y Grimm](#) subrayan la importancia de la comunicación de los modelos como una tarea esencial, lo que implica publicar o transferir el conocimiento producido a otras partes interesadas.

Por otro lado, [Fagiolo et al. \(2019\)](#), describen un enfoque algo similar, denominado calibración indirecta. Este enfoque también es iterativo y comienza con la identificación de hechos estilizados del mundo real que el modelador desea explicar. Luego sigue la especificación del

modelo, validación y prueba de hipótesis comparando la salida del modelo con observaciones del mundo real. Finalmente, podría haber un cuarto paso donde el modelo se emplea para análisis de políticas.

Es posible notar similitudes entre los enfoques de [Railsback y Grimm](#), y el de [Fagiolo et al.](#), ambos enfatizan la importancia de un proceso iterativo y la necesidad de validación y calibración. Sin embargo, [Railsback y Grimm](#) proporcionan una estructura más detallada para la formulación inicial del modelo, mientras que Fagiolo y su equipo enfatizan más en la validación comparando con datos del mundo real.

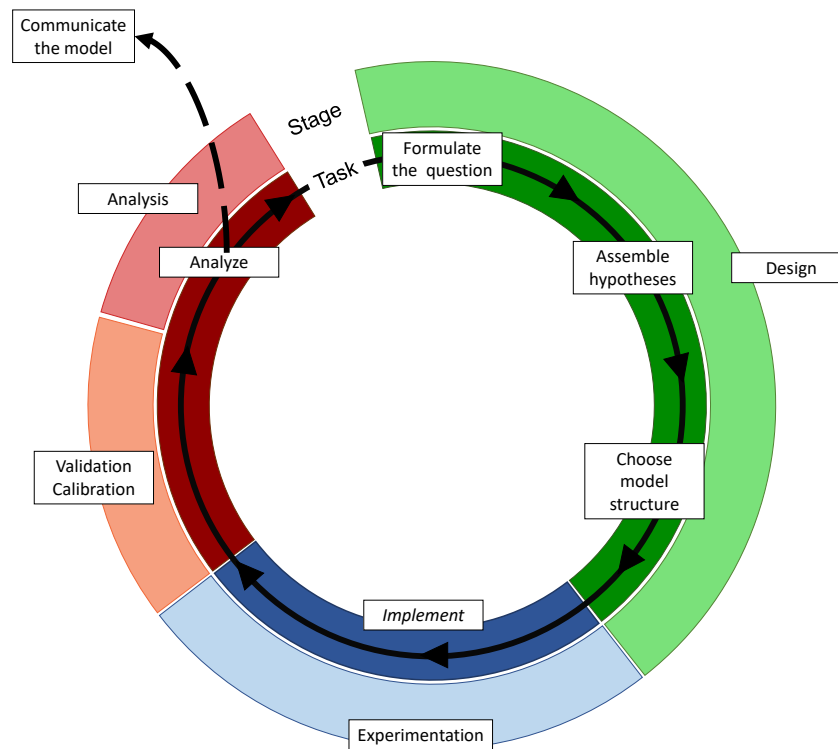
En la Figura 2.1, que representa el ciclo mejorado de la propuesta de [Railsback y Grimm \(2019\)](#), en el cual se destacan etapas y tareas. Este gráfico, alineado con la literatura revisada, segmenta el proceso en etapas de diseño, experimentación, validación-calibración y análisis. A diferencia de la propuesta original de [Railsback y Grimm \(2019\)](#), este ciclo mejorado incorpora etapas específicas y muestra cómo la validación y calibración son cruciales antes de pasar al análisis de la salida del modelo.

El modelado basado en agentes, debido a su naturaleza intrínsecamente compleja y multifacética, demanda un enfoque metodológico riguroso. Ambas propuestas presentadas por [Railsback y Grimm \(2019\)](#) y [Fagiolo et al. \(2019\)](#) enfatizan la necesidad de un proceso iterativo, que permita a los investigadores refinar y mejorar constantemente sus modelos para reflejar con precisión los sistemas que buscan representar. La iteratividad no solo garantiza la robustez del modelo, sino que también reconoce que nuestra comprensión del sistema subyacente es siempre una aproximación y puede ser perfeccionada. Lo más destacado es el reconocimiento universal de la validación y calibración como etapas vitales en este ciclo. Sin una validación adecuada, un modelo, por muy detallado o sofisticado que sea, corre el riesgo de desviarse significativamente de la realidad que busca emular. La representación gráfica del ciclo de desarrollo enriquece aún más esta discusión al visualizar la secuencia y la interrelación de las etapas involucradas en la construcción de un ABM efectivo.

2.1.5. Descripción y formulación

Habiendo establecido la importancia de un enfoque iterativo y la validación en el desarrollo de modelos basados en agentes, es imperativo abordar cómo se describen y formulan estos modelos de manera efectiva. Una descripción clara y una formulación precisa no solo facilitan

Figura. 2.1: Ciclo de desarrollo de un modelo basado en agentes.



Fuente: Platas-López et al. (2023b).

la comprensión y replicabilidad del modelo, sino que también establecen la base para una implementación efectiva.

Además, la reproducibilidad de los hallazgos en cualquier investigación es fundamental en el ámbito científico. En las comunidades de simulación, es esencial publicar modelos y métodos con precisión para promover el conocimiento y prevenir la redundancia. Esta cuestión es crucial, ya que los modelos pueden ser desarrollados y mantenidos por individuos o equipos, y donde los estudios que utilizan esos modelos pueden requerir auditorías o repeticiones.

Autores de diversas áreas científicas han explorado la reproducibilidad de modelos y han detectado que muchos informes de modelos publicados en revistas pueden ser ambiguos o incompletos, lo que dificulta su reutilización y extensión (Grimm et al., 2006; Dalle, 2012; Rahmandad y Sterman, 2012). Por lo tanto, este problema no es exclusivo de una disciplina en específico, lo que ha dado pie a llamados crecientes para crear directrices que ayuden a

los autores a informar completamente sobre sus modelos para maximizar la reproducibilidad y la transparencia.

En el contexto del Modelado Basado en Agentes, estas directrices se convierten en esenciales. La complejidad inherente a los ABM y la multiplicidad de formas en que pueden ser diseñados, desarrollados y calibrados subrayan la necesidad de estándares claros y robustos de reporte. Un ABM que no puede ser reproducido, evaluado o extendido por otros investigadores tiene un valor limitado en la comunidad científica y profesional.

En esta sección, se abordarán una variedad de métodos y pautas que han sido propuestos para la descripción y formulación de modelos y simulaciones. A medida que diversas disciplinas han buscado soluciones a los desafíos de la reproducibilidad, han surgido varios enfoques para abordar este problema. Estos métodos ofrecen estructuras y directrices específicas para documentar de manera exhaustiva y coherente los aspectos clave de un modelo, desde su concepción hasta su aplicación y resultados.

A pesar de la diversidad de propuestas que han emergido en esta área, el protocolo ODD (Overview, Design concepts, Details) ha demostrado ser uno de los enfoques más ampliamente aceptados y utilizados para la documentación de modelos basados en agentes. La metodología ODD ha demostrado ser un marco confiable y eficaz para describir modelos basados en agentes, emergiendo no solo en las ciencias ecológicas, sino también en diversas otras disciplinas. Por lo tanto, en esta tesis se adoptará el protocolo ODD como la metodología principal para documentar los modelos basados en agentes utilizados en el estudio, asegurando así una presentación rigurosa y clara de los modelos y su aplicación.

2.1.5.1. GASS

En su artículo, [Gass \(1984\)](#) destaca la importancia crucial de la documentación adecuada en los modelos basados en computadora, especialmente en contextos de toma de decisiones. Resalta que la falta de documentación es una razón significativa de la falla o utilidad deficiente de los modelos. Para abordar este desafío, [Gass](#) propone la creación de cuatro tipos de manuales para modelos computacionales a gran escala: el manual del analista, del usuario, del programador y del gerente. Estos manuales buscan consolidar información de diversas fases del ciclo de vida de un modelo.

El manual del analista está dirigido a las personas involucradas en el desarrollo, revisión y mantenimiento del modelo. Provee detalles técnicos necesarios para la comprensión y aplicación práctica del modelo. El manual del usuario tiene como objetivo ayudar a usuarios no programadores a comprender el propósito, capacidades y limitaciones del modelo. Explica la estructura general, los requisitos de datos de entrada, los formatos de salida y la interpretación de los resultados.

El manual del programador está orientado al personal de programación encargado del mantenimiento y modificación de los programas del modelo. Ofrece información detallada para depuración, corrección de errores, modificaciones y posibles transferencias de programas. El manual del gerente es esencial para ejecutivos que necesitan interpretar y utilizar los resultados del modelo, respaldar su uso continuo y gestionar su mantenimiento. Este manual proporciona una comprensión clara del propósito del modelo, sus capacidades, limitaciones, costos, beneficios y su rol en la organización.

[Gass](#) enfatiza que la documentación del modelo debe abarcar aspectos históricos, técnicos, de desarrollo, mantenimiento e implementación del modelo. El proceso de documentación debe ser concurrente con cada fase del ciclo de vida del modelo, siguiendo un enfoque estructurado y organizado. En última instancia, el autor sugiere que desarrollar una documentación integral es esencial para mejorar la usabilidad, confiabilidad y aceptación de modelos de decisión basados en computadora.

2.1.5.2. TRACE

Con el objetivo de mejorar la práctica actual de la modelización ecológica, [Schmolke et al. \(2010\)](#) proponen una estrategia en el contexto de la toma de decisiones ambientales. En su estudio, se resalta la importancia esencial de los modelos ecológicos en la exploración de las ramificaciones de políticas alternativas y escenarios de gestión.

No obstante, se destaca que la práctica actual de modelización adolece de ciertas deficiencias. A pesar de que se han identificado elementos de buenas prácticas en la literatura, estos a menudo son pasados por alto. El artículo sugiere que esto podría deberse a factores como la falta de participación de los tomadores de decisiones en el proceso de modelización, la carencia de incentivos para que los modeladores sigan buenas prácticas y la utilización de terminología inconsistente.

Estos autores proponen la adopción de un formato estándar para la documentación de modelos y sus respectivos análisis, denominado “Transparent and Comprehensive Ecological Modeling” (TRACE). El propósito de este formato es lograr una mayor transparencia y coherencia en la documentación de todo el proceso de modelización, desde el desarrollo inicial hasta la aplicación práctica en la toma de decisiones. La estructura de la documentación TRACE abarca tres etapas fundamentales.

En primer lugar, se encuentra el “Desarrollo del modelo”, que comprende la formulación del problema, el contexto de uso, el público objetivo y las preguntas que el modelo busca responder. Se establece su aplicabilidad y se evalúa la disponibilidad de datos y conocimientos necesarios. Se describen detalladamente la concepción y formulación del modelo, incluyendo complejidad y suposiciones. La implementación se detalla, abarcando programas, plataformas y scripts. La parametrización incluye valores de parámetros, fuentes de datos e incertidumbres. La calibración documenta conjuntos de datos, parámetros y métodos de optimización.

La siguiente etapa es la de “Pruebas y análisis del modelo”. Aquí, se verifica si el modelo opera según sus especificaciones y se documentan las pruebas realizadas. Se explora la sensibilidad del modelo mediante variación de parámetros, justificando combinaciones y rangos. La validación compara salidas del modelo con datos empíricos no usados en la calibración, indicando fuentes y métodos.

Por último, se aborda la “Aplicación del modelo”, donde se presentan resultados relevantes para decisiones. Describe experimentos de simulación, estadísticas para analizar salidas y aborda incertidumbre en estas, incluyendo variabilidad, ruido, sesgo y estocasticidad. Detalla cómo evaluar la incertidumbre a través de diferentes modelos o submodelos. Las recomendaciones finales se derivan de las salidas, especificando extrapolaciones temporales y espaciales.

Los autores enfatizan que el enfoque TRACE no solo sirve como una herramienta para los modeladores, sino que también facilita la evaluación de los modelos por parte de los tomadores de decisiones y otras partes interesadas. Proporciona una guía estructurada y comprensiva para documentar todos los aspectos del proceso de modelización, lo que a su vez fomenta una mayor transparencia y coherencia en el uso de modelos en la toma de decisiones ambientales.

2.1.5.3. MIASE

Unas directrices fundamentales para la simulación de procesos biológicos fueron desarrolladas por [Waltemath et al. \(2011\)](#), las cuales denominaron “Minimum Information About a Simulation Experiment” (MIASE). Estas pautas han demostrado ser una herramienta útil para habilitar la reutilización del trabajo existente en la biología moderna. Mientras que las “Minimum Information Required in the Annotation of Models” (MIRIAM) promueven el intercambio y la reutilización de modelos computacionales bioquímicos, es importante destacar que la información sobre un modelo por sí sola no es suficiente para permitir su reutilización eficiente en un entorno computacional. Los algoritmos numéricos avanzados y los flujos de trabajo de modelado complejos utilizados en la biología computacional moderna dificultan la reproducción de las simulaciones. Por lo tanto, es esencial definir la información central necesaria para realizar simulaciones de esos modelos.

MIASE establece el conjunto mínimo de información que debe proporcionarse para hacer que la descripción de un experimento de simulación esté disponible para otros. Esto incluye la lista de modelos a utilizar y sus modificaciones, todos los procedimientos de simulación a aplicar y en qué orden, el procesamiento de los resultados numéricos en bruto y la descripción de la salida final. MIASE permite la reproducción de cualquier experimento de simulación. La provisión de esta información, junto con un conjunto de modelos requeridos, garantiza que el experimento de simulación represente la intención de los autores originales. Siguiendo las directrices de MIASE, se mejorará la calidad de los informes científicos y también se permitirán esfuerzos colaborativos y más distribuidos en el modelado y la simulación computacional de procesos biológicos.

El alcance de MIASE se divide en tres categorías. El primer nivel de reglas enumera la información que debe proporcionarse sobre los modelos que se utilizarán en el experimento de simulación. Todos los modelos deben ser enumerados o descritos de manera que permita la reproducción del experimento. Las reglas de segundo nivel especifican cómo describir el experimento de simulación en sí. Toda la información necesaria para ejecutar cualquier paso del experimento debe ser proporcionada. Finalmente, las reglas del tercer nivel tratan sobre los resultados devueltos por el experimento. Una publicación que describa un experimento de simulación debe cumplir con los tres niveles de reglas para que la descripción se considere compatible con MIASE.

Las directrices de MIASE son un paso esencial hacia la mejora de la calidad de la investigación científica al garantizar la reutilización y la reproducción confiable de los experimentos de simulación. Estas pautas proporcionan una estructura coherente y unificada para la descripción y formulación de experimentos de simulación en biología computacional, permitiendo así una mayor transparencia, colaboración y desarrollo de modelos en esta área.

2.1.5.4. MMRR

Rahmandad y Sterman (2012) desarrollaron los Minimum Model Reporting Requirements (MMRR), directrices de presentación para la investigación basada en simulaciones en las ciencias sociales, enfocándose específicamente en escenarios de investigación de dinámica de sistemas. Los autores resaltaron la importancia de la reproducibilidad en la investigación basada en simulaciones y la falta de directrices estandarizadas de presentación en este ámbito. Enfatizaron que la reproducibilidad es esencial para verificar los resultados de la investigación, construir sobre trabajos anteriores y establecer la confiabilidad de los hallazgos.

Para demostrar la necesidad de una presentación mejorada, Rahmandad y Sterman (2012) llevaron a cabo un análisis de los artículos publicados en la revista *System Dynamics Review* entre 2010 y 2011. Entre los artículos que informaban resultados de simulación, encontraron que muchos carecían de documentación suficiente de las ecuaciones del modelo, los valores de los parámetros y las entradas exógenas. Esta deficiencia en la presentación obstaculizó la reproducibilidad de los resultados y destacó la necesidad de directrices exhaustivas de presentación.

Las directrices de presentación propuestas consisten en requisitos mínimos y preferidos para diferentes aspectos de la presentación de la investigación basada en simulaciones. Los autores recomiendan prácticas claras de visualización, como evitar la saturación en diagramas de relaciones causales y utilizar convenciones de nomenclatura estándar para las variables. También sugieren el uso de diversas herramientas de visualización, como diagramas de subsistemas y diagramas de bucles causales.

Los MMRR exigen la presentación clara de las operaciones computacionales, las ecuaciones, las reglas algorítmicas, los valores de los parámetros y las condiciones iniciales. También se requieren unidades de medida para las variables y parámetros. Introdujeron PMRR como

una mejora, que incluye proporcionar fuentes de datos, definiciones de variables y código fuente para la implementación del modelo.

Para la presentación de experimentos de simulación, los requisitos mínimos (MSRR) abarcan descripciones detalladas de las configuraciones de simulación, las plataformas de software y hardware, los algoritmos de simulación, los ajustes de parámetros, el número de iteraciones y los pasos de post-procesamiento. La presentación preferida de simulación (PSRR) incluye información adicional sobre el análisis de sensibilidad, los costos computacionales, las medidas de incertidumbre y la significación estadística.

En cuanto a la presentación de experimentos de optimización, MORR implica especificar las funciones objetivo de optimización, los algoritmos de búsqueda, los espacios de parámetros y otros detalles relevantes. Los autores introdujeron PORR como una extensión, recomendando la inclusión de códigos de implementación de optimización, costos computacionales, medidas de incertidumbre para parámetros estimados y medidas de confianza en los resultados de optimización.

Los autores enfatizaron la importancia de estas directrices para facilitar la colaboración, la construcción de confianza y el progreso general de la investigación basada en simulaciones. Reconocieron que las herramientas de software, los requisitos de los asesores y las políticas de revistas pueden desempeñar un papel crucial en la promoción de la adopción de estas directrices en toda la comunidad de investigadores.

2.1.5.5. CHEERS

En el año 2013, [Husereau et al.](#) abordaron la complejidad de la presentación de evaluaciones económicas de intervenciones en salud mediante la creación del Consolidated Health Economic Evaluation Reporting Standards (CHEERS) statement. La elaboración de este conjunto de pautas buscaba solventar los desafíos inherentes a la comunicación de resultados económicos en el campo de la salud y proporcionar una guía actualizada y amigable para los investigadores, editores y revisores.

En el centro de la motivación para el desarrollo del CHEERS statement se encontraba la necesidad de mejorar la transparencia y claridad en la presentación de evaluaciones económicas en salud. A diferencia de los estudios clínicos, las evaluaciones económicas requieren

espacio adicional para elementos como el uso de recursos, los costos y los resultados relacionados con las preferencias, lo cual complica la presentación efectiva. Además, la creciente cantidad de publicaciones y el potencial impacto de decisiones basadas en hallazgos mal comunicados subrayaron la importancia de promover estándares de presentación sólidos.

La metodología utilizada para desarrollar el CHEERS statement fue rigurosa y multidisciplinaria. Se llevaron a cabo rondas de Delphi modificadas en las que participaron expertos de diversas áreas, incluyendo academia, práctica clínica, industria, gobierno y edición. A través de este proceso, se identificaron 24 elementos esenciales para la presentación de evaluaciones económicas en seis categorías principales: título y resumen, introducción, métodos, resultados, discusión y otros.

La contribución más significativa del CHEERS statement radica en su capacidad para facilitar la interpretación y comparación de estudios económicos en salud. Al proporcionar una lista de verificación de elementos esenciales, el CHEERS busca optimizar la calidad de la presentación de resultados, asegurando que los detalles críticos, como los métodos de evaluación, los resultados económicos y las consideraciones de costo-efectividad, sean comunicados de manera coherente y completa.

La importancia del CHEERS statement se destaca aún más por su enfoque en la neutralidad metodológica. Aunque establece criterios mínimos para la presentación, permite a los analistas de evaluaciones económicas la libertad de elegir métodos diferentes según el contexto específico. Además, el enfoque en la multidisciplinariedad y la colaboración internacional respalda su aplicabilidad en una variedad de escenarios y sistemas de salud.

2.1.5.6. STRESS

Al igual que en otros ámbitos, en la investigación de [Monks et al. \(2018\)](#), se aborda una cuestión que es fundamental en la Investigación Operativa y la Ciencia de la Gestión: la reproducibilidad de los resultados en los estudios de simulación. Los autores enfatizan que la falta de detalles completos y claros en la presentación de modelos de simulación puede tener efectos perjudiciales en la capacidad de los investigadores para comprender, reutilizar y ampliar los resultados de simulación en diferentes métodos, como la dinámica de sistemas, la simulación por eventos discretos y la simulación basada en agentes.

Para abordar esta problemática, los autores proponen las pautas “STRESS” (Strengthening The Reporting of Empirical Simulation Studies), que se dividen en tres listas de verificación: STRESS-DES, STRESS-ABS y STRESS-SD, correspondientes a simulación por eventos discretos, simulación basada en agentes y dinámica de sistemas, respectivamente. Estas listas de verificación contienen 20 elementos cada una y están diseñadas para guiar a los investigadores en la presentación efectiva de modelos de simulación y resultados.

Las pautas STRESS se dividen en seis secciones clave: objetivos, lógica del modelo, datos, experimentación, implementación y acceso al código. En la sección de objetivos, se insta a los autores a reportar claramente los propósitos y el marco experimental del estudio. En la lógica del modelo, se enfatiza la importancia de describir el modelo base y las diferencias en la lógica entre diferentes escenarios. Además, se sugiere el uso de diagramas de simulación reconocidos para comunicar el diseño del modelo. En la sección de datos, se subraya la necesidad de detallar los parámetros de entrada, incluidas fuentes de datos, preprocesamiento y suposiciones. La sección de experimentación aborda cómo se inicia el modelo, la longitud de ejecución y el enfoque de estimación de resultados. La implementación se refiere a los detalles técnicos, como el software utilizado y la especificación del sistema. Finalmente, se recomienda proporcionar acceso al código del modelo para otros investigadores.

Las pautas STRESS no solo son relevantes para académicos, sino también para profesionales de la industria. Seguir estas pautas puede aumentar la reproducibilidad de los resultados y facilitar la revisión por parte de los pares y los editores. Si bien estas pautas brindan un marco sólido, se reconoce que algunas situaciones pueden requerir adaptaciones y flexibilidad en su aplicación.

2.1.5.7. RAT-RS

[Achter et al. \(2022\)](#) proporcionan un marco integral para mejorar la documentación del uso de datos en la Modelización Basada en Agentes. Los autores destacan la necesidad de estándares de informes que aborden no solo el desarrollo de modelos, sino también el aspecto crítico de la utilización de datos, que a menudo ha sido pasado por alto en la investigación del ABM. Los autores enfatizan que un enfoque estándar de informes debe abarcar diversos usos de los datos, como la especificación, calibración y validación, y debe

ser compatible con diferentes tipos de datos, incluidos los datos de encuestas, etnográficos y experimentales.

Los autores argumentan que el establecimiento de estándares de informes es esencial para el avance de la investigación, fomentando la transparencia y permitiendo una mejor evaluación y replicación de los ABM. Argumentan que las disciplinas científicas exitosas establecen estándares que guían las prácticas de investigación y los criterios de evaluación. Sin embargo, debido a la naturaleza creativa e iterativa del desarrollo de los ABM, los autores afirman que la documentación del uso de datos se vuelve desafiante. A menudo, los investigadores tienen dificultades para reconstruir cómo y por qué se utilizaron los datos en ABM específicos, lo que lleva a dificultades en la evaluación y replicación de los resultados de la investigación.

Los autores proponen una solución al introducir el RAT-RS (Rigour And Transparency – Reporting Standard), un conjunto de preguntas diseñadas para documentar de manera exhaustiva el uso de datos en ABM. Enfatizan que el enfoque RAT-RS tiene como objetivo proporcionar una comprensión clara de las motivaciones detrás del uso de datos, facilitar la evaluación y permitir la replicación. Los autores argumentan que el enfoque RAT-RS se diferencia de los estándares de documentación existentes en su enfoque en la motivación del uso de datos, su generalidad en todas las disciplinas y su evitación de suposiciones sobre los tipos de datos o metodologías preferidas.

En cuanto a la implementación práctica, los autores enfatizan que el RAT-RS está diseñado para ser fácil de usar y no debe imponer perspectivas rígidas a los investigadores. Recomendamos completar las preguntas del RAT-RS durante el proceso de modelado para garantizar que las decisiones sobre el uso de datos se capturen adecuadamente. Sin embargo, los autores reconocen que el RAT-RS también se puede utilizar retrospectivamente para resumir la utilización de datos después de la finalización de la investigación.

En general, los autores sostienen que el RAT-RS ofrece una contribución valiosa a la comunidad de ABM al abordar los desafíos de documentar el uso de datos y promover la transparencia y el rigor en un campo de investigación en desarrollo. Proporcionan información sobre el desarrollo del RAT-RS a través de un proceso colaborativo e iterativo que involucra talleres, pruebas y comentarios de un grupo diverso de investigadores. Los autores concluyen destacando los beneficios potenciales del RAT-RS, incluida la mejora de las evaluaciones de

modelos, el conocimiento procedimental compartido y el progreso científico mejorado dentro del ámbito del ABM.

2.1.5.8. ODD

Antes de la aparición de TRACE ([Schmolke et al., 2010](#)), el protocolo ODD (Overview, Design concepts, Details) se destacó como un pionero en el campo de la documentación de modelos basados en agentes en el ámbito ecológico. Inicialmente presentado en 2006 por [Grimm et al. \(2006\)](#), ODD es un documento estandarizado que proporciona una descripción coherente, lógica y legible de la estructura y la dinámica de los ABM.

Aunque la primera versión de ODD dejó algunos aspectos sin abordar, su actualización en 2010 ([Grimm et al., 2010](#)) superó algunas de estas limitaciones. A partir de esa actualización, ODD ([Grimm et al., 2010, 2020](#)), consta de siete elementos como se muestra en la Figura 2.2. Estos elementos se agrupan en tres niveles de descripción, que van de lo general a lo específico.

Overview. Una descripción general del modelo, incluyendo su propósito y sus componentes básicos: agentes, variables que los describen y el entorno, y escalas utilizadas en el modelo, por ejemplo, tiempo y espacio; así como una visión general de los procesos y su programación.

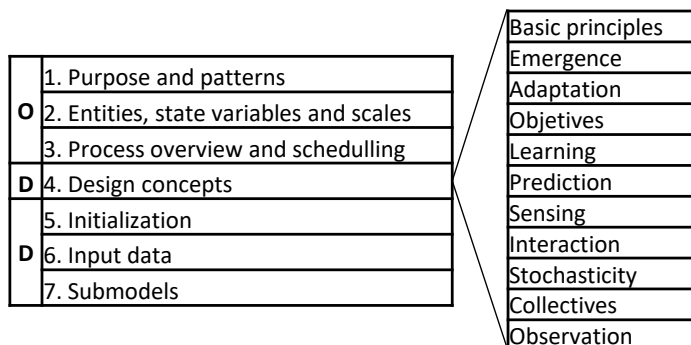
Design concepts. Una breve descripción de los principios básicos que subyacen al diseño del modelo, por ejemplo, racionalidad, emergencia, adaptación, aprendizaje, etc.

Details. Definiciones completas de los submodelos involucrados.

Además, existen 11 conceptos de diseño específicos, si alguno de ellos no aplica, se pueden omitir, y si un modelo incluye un concepto importante de interés general, que aún no forma parte de ODD, se puede agregar al final de la sección. 4 “Conceptos de diseño”. Dependiendo de la tarea a resolver a través de ML en ABM, se recomienda incluir la descripción de la implementación en esa sección.

A lo largo de los años, ODD ha evolucionado para convertirse en un método ampliamente aceptado y utilizado en la documentación de ABM. Su utilidad no se limita únicamente a las ciencias ecológicas, sino que ha sido aplicado exitosamente en diversas disciplinas. Además, ha continuado mejorando con el tiempo, y su versión actualizada en 2020 ([Grimm et al., 2020](#))

Figura. 2.2: Estructura de descripciones de modelos siguiendo el protocolo ODD.



Fuente: Adaptado de Grimm et al. (2010, 2020).

lo ha consolidado aún más como el enfoque preferido para brindar una comprensión sólida y detallada de los modelos basados en agentes.

2.1.6. Críticas a los ABM

A medida que los modelos basados en agentes han emergido como una herramienta poderosa para simular y analizar sistemas complejos en diversas disciplinas, también han surgido críticas y preocupaciones sobre su aplicación y fiabilidad. Estas críticas abarcan desde consideraciones prácticas, como el costo computacional y la curva de aprendizaje requerida, hasta retos más fundamentales, como los desafíos inherentes en el diseño y calibración del modelo. Estas objeciones no solo provienen de detractores del enfoque ABM, sino también de académicos que, a pesar de reconocer sus potenciales beneficios, han señalado áreas de mejora y precaución.

A través de un exhaustivo análisis en plataformas académicas prominentes como Google Scholar, Web of Science y Scopus, se ha reunido una gama de críticas hacia los ABM, abarcando un periodo desde el año 2000 hasta la fecha actual, un tiempo durante el cual el interés en los ABM ha crecido significativamente, como se aprecia en la Figura 1.1. Esta subsección busca sintetizar estas críticas, proporcionando una perspectiva equilibrada sobre los retos que los modelos basados en agentes enfrentan, y estableciendo un fundamento sólido para abordarlos y superarlos en investigaciones futuras. Estas críticas, que han sido sintetizadas y resumidas para una fácil referencia, se pueden consultar en la Tabla 1.1 del capítulo 1.

En este trabajo, es crucial reconocer que no todas las críticas a los ABM serán objeto de estudio. Específicamente, no serán abordadas las críticas relacionadas con la curva de aprendizaje en programación y el costo computacional. En esta tesis, nos concentramos en las críticas asociadas al diseño, calibración de parámetros e interpretación, ya que identificamos que el aprendizaje automático ofrece un potencial significativo para abordar estas áreas de preocupación.

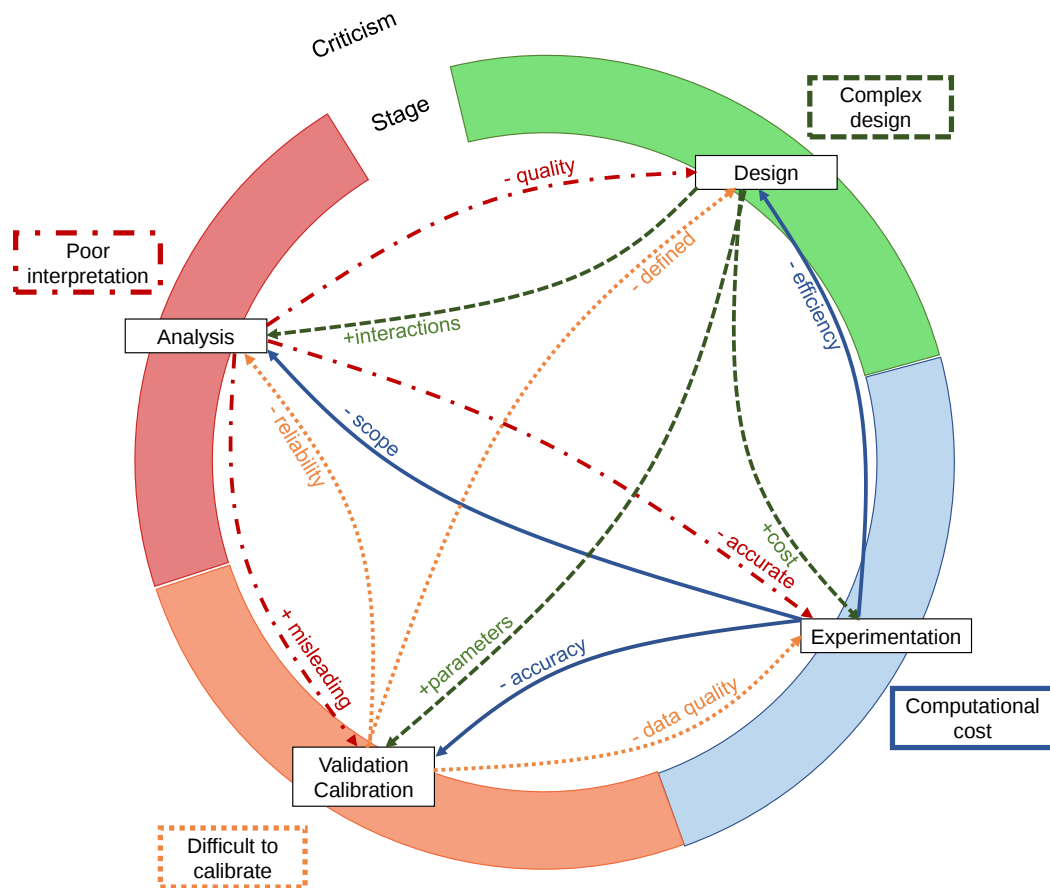
La interacción y propagación de críticas a través de las etapas del desarrollo de un ABM es compleja y requiere un examen cuidadoso. A modo de ilustración:

- Un diseño intrincado y complejo en la etapa inicial conlleva múltiples repercusiones. Por un lado, este diseño incrementa el costo computacional, introduce más parámetros que necesitan calibración y validación, y añade múltiples interacciones que complican el análisis posterior.
- La existencia de un elevado costo computacional puede, a su vez, afectar la eficiencia en el diseño, reducir la precisión en la calibración y validación, y limitar el alcance del análisis.
- Una inadecuada calibración o validación repercute en la calidad de los datos utilizados para la experimentación, compromete la definición del diseño y cuestiona la fiabilidad de cualquier análisis posterior.
- Finalmente, una interpretación incorrecta o sesgada puede llevar a errores en la validación y calibración, impactar negativamente la calidad de la experimentación, y comprometer la integridad del diseño original.

Estas interacciones, más que aisladas, son cascadas de efectos que afectan la integridad del modelo y, como [Siegfried \(2014\)](#) destacó, una decisión equivocada en una etapa puede propagar errores y complicaciones en las etapas subsiguientes, llevando a modelos de alto consumo computacional y calibraciones incorrectas. La Figura 2.3 ilustra detalladamente estas relaciones y cómo se entrelazan las críticas principales hacia los ABM con las etapas de su desarrollo.

En las subsecciones que se presentan a continuación, profundizaremos en cada una de estas críticas, ofreciendo una descripción detallada y un análisis pormenorizado de sus implicaciones y alcances en el contexto del modelado basado en agentes.

Figura. 2.3: Relación sobre las principales críticas hacia los modelos basados en agentes y su correspondiente etapa de modelado.



Fuente: Platas-López et al. (2023b).

2.1.6.1. Computacionalmente caro

Si bien las ecuaciones matemáticas proporcionan teoremas sólidos tras una sola ejecución, [Axtell \(2000\)](#) cuestiona la robustez de los teoremas producidos por los ABM. Pero, ¿qué significa realmente esta falta de robustez? ¿Es simplemente una limitación intrínseca de los ABM o una oportunidad para una inmersión más profunda en la simulación?

Al abordar esta pregunta, [Charteris et al. \(2001\)](#) refuerza la perspectiva de [Axtell](#) al señalar la inconsistencia que puede surgir cuando se simulan múltiples agentes con comportamientos heterogéneos. Estas inconsistencias no son solo simples variaciones; pueden ser tendencias o estructuras enteras que aparecen y desaparecen entre diferentes iteraciones. Si bien esto puede sonar alarmante, hay un hilo común entre [Charteris et al.](#) y [Axtell](#): la necesidad de múltiples ejecuciones para evaluar la robustez.

A medida que nos sumergimos más en la problemática de los ABM, surge una preocupación técnica: la limitante de capacidad. [Breckling \(2002\)](#) toca este punto señalando que solo un número finito de agentes puede ser representado en un ABM. ¿Pero es realmente esto una restricción o simplemente una invitación a ser más selectivo y estratégico en nuestra modelización?

Aunque los desafíos son claros, también hay un impulso hacia la integración y la adaptación. [Nourqolipour y Shariff \(2010\)](#) sugiere que al combinar ABM con herramientas de manejo de datos como GIS se puede superar su limitación en la representación de patrones fenoménicos. De manera similar, [Rand y Rust \(2011\)](#) valoran la capacidad del ABM para modelar con mayor realismo y complejidad, mientras que [Conte y Paolucci \(2014\)](#) y [Kostadinov et al. \(2014\)](#) ven una oportunidad para ABM en la era del “Big Data”.

Dicho esto, no podemos ignorar la crítica contundente de [Wilensky y Rand \(2015\)](#): las ecuaciones pueden ser más compactas y rápidas. Pero como contrapunto, [Manzo \(2014\)](#) advierte sobre la idealización del modelado basado en ecuaciones, insinuando que no es inherentemente superior a los ABM.

En este diálogo complejo y en constante evolución, [Tarvid \(2016\)](#) nos brinda una perspectiva equilibrada al reconocer que, aunque el ABM ofrece modelos más realistas, todavía estamos operando en el reino de la modelización y no de la duplicación de la realidad, por lo

que nos invita a ser selectivos en los elementos que pueden abstraer mejor el fenómeno bajo estudio.

La advertencia de [Broniec et al. \(2021\)](#) sobre la complejidad temporal de los ABM solo solidifica la idea central de este diálogo: los ABM son poderosos, pero su poder viene con desafíos inherentes que los investigadores deben abordar con cuidado y estrategia, con el objetivo de evitar una complejidad innecesaria en nuestros modelos.

2.1.6.2. Diseño

El diseño de modelos basados en agentes es una tarea que, si bien ofrece potenciales ventajas, se encuentra empañada por distintos desafíos. A primera vista, parece que la mayor representación de un sistema es beneficiosa. Sin embargo, [Charteris et al. \(2001\)](#) advierten que el deseo de una representación más detallada puede conducir a modelos altamente parametrizados, empujando a los investigadores fuera de su área de especialidad. Esta es una preocupación compartida por varios investigadores.

Por otro lado, mientras que [Breckling \(2002\)](#) aborda la versatilidad de los ABM en conectar diferentes niveles jerárquicos, también reconoce las limitaciones inherentes en tender puentes entre espectros de diferente escala. Pero, ¿es esto realmente una crítica al diseño per se, o más bien un reflejo de nuestra incapacidad actual para conectar niveles jerárquicos? [Ligmann y Sun \(2010\)](#) complementan esta discusión al resaltar cómo la flexibilidad en la representación de comportamientos, aunque ventajosa, introduce incertidumbres en el diseño.

En el contexto macroeconómico, [Lengnick \(2013\)](#) observa que la naturaleza flexible de los ABM da a los modeladores una vasta libertad, lo que podría ser tanto una fortaleza como una debilidad. Mientras que en el modelado convencional existen límites naturales, los modeladores de ABM tienen que autodisciplinarse para evitar la sobrecomplejización. Coincido con [Lengnick](#) en que esta disciplina es fundamental, pero no es inherentemente un defecto del ABM. En cambio, [Conte y Paolucci \(2014\)](#) añaden otra dimensión a este debate al señalar la falta de consistencia y fundamentación teórica en muchos modelos de agentes. Aquí, pareciera que el núcleo del problema es la falta de un protocolo estandarizado para la implementación de modelos.

[Ghorbani et al. \(2014\)](#) y [Siegfried](#) presentan desafíos relacionados con la simulación del

entorno y la elección del nivel de detalle. Ambas críticas tocan un punto crucial: el balance entre precisión y comprensibilidad. Por último, [Tarvid \(2016\)](#) apunta hacia un problema más fundamental, la comprensión de los mecanismos del comportamiento humano. ¿Podemos realmente modelar lo que no entendemos completamente?

En resumen, mientras los modelos basados en agentes ofrecen una herramienta potente y flexible, es evidente que esta misma flexibilidad puede ser una espada de doble filo. La disciplina, la coherencia y la humildad parecen ser esenciales en la construcción de estos modelos.

2.1.6.3. Calibración

La calibración en los modelos basados en agentes es un desafío trascendental y recurrente en la literatura. Es un campo de batalla donde la parametrización se encuentra con la precisión y la complejidad.

En primer lugar, [Breckling \(2002\)](#) arroja luz sobre una problemática intrínseca de la calibración: la parametrización. Según él, en los ABM no basta con generalizaciones y abstracciones burdas. Pero, ¿cuál es el motivo detrás de esta necesidad de precisión? Al profundizar, nos damos cuenta de que, mientras en estos modelos muchos parámetros son insensibles, a nivel multi-agente, funciones enteras pueden ser tremendamente sensibles a cambios diminutos. La pregunta que surge es: ¿cómo equilibramos la precisión sin caer en la trampa de la sobre-parametrización?

Aquí es donde [Ligmann y Sun \(2010\)](#) entra en escena con un enfoque más funcional. Plantea una crítica válida: a menudo, la exploración ABM carece de indicadores claros. Sin estos, nos encontramos navegando a ciegas, incapaces de discernir la influencia real de las entradas en los resultados del modelo. [Ligmann y Sun](#) parece sugerir que necesitamos herramientas más claras para descifrar la incertidumbre y destacar las regiones críticas del modelo. Entonces, ¿es la falta de herramientas analíticas lo que nos impide una calibración exitosa?

Sin embargo, [Siegfried](#) aporta otra dimensión al debate, centrándose en el nivel de detalle de los modelos. Argumenta que un nivel de detalle incorrecto puede ser nuestra perdición, llevándonos a modelos excesivamente complicados que requieran una calibración infinita.

Esta perspectiva resuena con lo que [Breckling \(2002\)](#) mencionaba sobre la precisión, pero añade un matiz importante: la sencillez y la precisión no son mutuamente excluyentes.

Entonces, al cruzar estos puntos de vista, emerge una imagen más compleja de la calibración en ABM. Parece que el camino hacia una calibración efectiva es tenue, requiere tanto de precisión como de la correcta selección de detalles, complementada con herramientas analíticas robustas.

2.1.6.4. Interpretación

Los desafíos en la interpretación de los Modelos Basados en Agentes en las ciencias sociales son innegables. Uno podría pensar que la naturaleza compleja y, a menudo, impredecible de los patrones y resultados de interacciones entre agentes es un reto constante, como afirma [Li et al. \(2008\)](#). Esta impredecibilidad inherente a las interacciones de los agentes también resuena con los argumentos de [Siegfried \(2014\)](#), quien indica la dificultad de prever cómo el comportamiento a nivel micro podría afectar el comportamiento del modelo en general. Es decir, hay una tendencia emergente que complica la interpretación.

Sin embargo, esta complejidad en la interpretación no se limita únicamente a la emergencia. [Lengnick \(2013\)](#) va un paso más allá al sugerir que, debido a la flexibilidad inherente en el diseño de modelos macroeconómicos, los modelos ABM varían en naturaleza y complejidad, complicando aún más su interpretación. De hecho, es esta misma variedad lo que dificulta determinar qué patrón macro es el resultado de qué micropropiedad. A esto, [Ghorbani et al. \(2014\)](#) agregan que la emergencia en simulaciones sociales, por ser visualmente no reconocible, imposibilita su análisis y eventual control.

A pesar de estos desafíos, se ha reconocido que el problema real surge cuando se intenta validar estos modelos. [Baldwin et al. \(2015\)](#) y [Lee et al. \(2015\)](#) comparten esta perspectiva, con el último detallando varios desafíos que rodean el análisis de resultados de un ABM, como la visualización y problemas estadísticos. La comunicación efectiva de estos modelos a audiencias no técnicas también se destaca como una preocupación, lo que muestra la amplitud de las limitaciones al interpretar y presentar los resultados de los ABM.

Finalmente, [Tarvid \(2016\)](#) ofrece una perspectiva ligeramente diferente, apuntando a la juventud del método ABM como una fuente de dificultad. La falta de un flujo de trabajo

estándar y la incomparabilidad de los modelos debido a su contenido teórico heterogéneo resaltan la necesidad de una mayor estandarización en el campo.

2.1.6.5. Necesidad de programación

La complejidad inherente al diseño de Modelos Basados en Agentes y el alto grado de habilidades de programación requerido ha sido un tema recurrente en la literatura. Por ejemplo, [Tsfatsion \(2002\)](#) apuntó hacia una barrera de entrada para aquellos economistas y otros científicos no familiarizados con la programación, resaltando la inclinación de herramientas como C++ y Java para programadores experimentados. Esta perspectiva fue corroborada por [Baldwin et al. \(2015\)](#), quienes, al contrastar el enfoque de modelación basado en eventos con un ABM, identificaron la capacidad de programación como una limitante potencial.

Por otro lado, mientras que algunos argumentan sobre los límites intrínsecos de los ABM para representar la vasta complejidad del comportamiento humano, [Malleson et al. \(2010\)](#) proponen que dichas limitaciones no son intrínsecas a los ABM per se, sino que se originan en el nivel de detalle que uno elige para el modelo. Es curioso cómo este argumento resuena con [Siegfried](#), quien lamenta la dificultad de discernir el nivel micro del comportamiento que produce comportamientos deseados a nivel macro. Estos argumentos convergen en la esencia de un dilema: ¿Hasta qué punto deberíamos simplificar o complicar un modelo para que sea tanto manejable como representativo?

La convergencia entre [Tsfatsion \(2002\)](#) y [Baldwin et al. \(2015\)](#) sugiere una interrogante fundamental: ¿Está el desafío realmente en la complejidad del ABM o en nuestra capacidad (o falta de ella) para programar adecuadamente dicho modelo? Y aún más, si uno acepta la premisa de [Malleson et al.](#) y [Siegfried](#) sobre la elección del nivel de detalle, entonces es crucial que las herramientas y lenguajes que utilicemos sean lo suficientemente versátiles para permitir esta flexibilidad sin aumentar indebidamente la barrera de entrada para los novatos.

En definitiva, un equilibrio debe ser encontrado entre la accesibilidad y la precisión en el modelado basado en agentes. Mientras que las herramientas y lenguajes actuales pueden favorecer a aquellos con habilidades de programación más avanzadas, el campo debe moverse hacia soluciones más inclusivas, especialmente si quiere atraer a un público más amplio y diverso.

2.2. Aprendizaje Automático

El avance tecnológico y la acumulación masiva de datos en la era digital han impulsado la necesidad de comprender y aplicar técnicas que permitan a las máquinas aprender de manera autónoma, dando origen al campo del aprendizaje automático. Esta sección se sumerge en el corazón de esta disciplina, explorando sus fundamentos y delineando su relevancia y aplicabilidad en diversos contextos.

Inicialmente, se abordará la esencia del aprendizaje desde una perspectiva general, profundizando en su conceptualización dentro del ámbito computacional. Esta primera parte provee una base sólida para entender cómo las máquinas pueden ser capacitadas para aprender y adaptarse a través de la experiencia y la información.

Posteriormente, se presentará una detallada exploración de los diversos tipos de aprendizaje, desde técnicas tradicionales hasta las más avanzadas. Más allá de una simple categorización, se buscará identificar las ventajas intrínsecas de cada enfoque y discernir la medida en que estos pueden ser implementados para potenciar y mejorar los Modelos Basados en Agentes.

Con esta estructura, la sección aspira a ofrecer al lector una comprensión clara y concisa del aprendizaje automático, así como una visión crítica de su potencial y limitaciones en el contexto de los ABM.

2.2.1. Aprendizaje

El aprendizaje, en palabras de [Michalski et al. \(1983\)](#), es fundamental para cualquier comportamiento inteligente. De hecho, cualquier variación duradera en el comportamiento que emerge de una experiencia anterior se define ampliamente como aprendizaje. Esta habilidad de aprender se manifiesta en diversas formas, desde lo más trivial hasta aspectos profundos y complejos. Cuando una computadora es el agente que aprende, nos referimos a esto como aprendizaje automático. En términos básicos, una computadora examina datos, crea un modelo basado en esos datos y utiliza dicho modelo como una hipótesis acerca del mundo y un software diseñado para resolver problemas, detectando patrones en los datos, según [\(Russell et al., 2020\)](#).

Este ámbito investiga las maneras en que las computadoras pueden mejorar con base

en experiencias anteriores, situándose en el cruce entre la computación y la estadística. Es más, todos los agentes, sin excepción, pueden potenciar su desempeño mediante el aprendizaje automático. Cualquier clase de agente tiene el potencial de transformarse en un agente de aprendizaje. La idea de crear sistemas inteligentes que puedan aprender no es reciente. Alan Turing ya lo había propuesto en 1950, sugiriendo que se construyeran máquinas capaces de aprender y, posteriormente, enseñarles. Esta propuesta surgió como alternativa al enfoque de programar máquinas inteligentes manualmente, un proceso arduo y complicado, especialmente cuando no se tiene una comprensión completa del funcionamiento de los entornos.

Así, cuando se busca diseñar un sistema inteligente, surge la necesidad de discernir entre las fuentes de inteligencia que deben programarse y aquellas que pueden o deben ser adquiridas a través del aprendizaje. [Russell et al. \(2020\)](#) destaca el aprendizaje como el medio para construir sistemas capaces y para expandir la capacidad del diseñador a territorios desconocidos.

En ocasiones, los diseñadores tienen un modelo sólido del agente y su entorno. Sin embargo, muchas veces este no es el caso, obligando al agente a confiar en datos de experiencias pasadas para tomar decisiones. [Russell et al. \(2020\)](#) diferencia entre entornos conocidos y desconocidos, no basándose en el entorno per se, sino en el conocimiento que el agente o diseñador tiene sobre las “leyes físicas” de dicho entorno. Evidentemente, en ambientes desconocidos, es vital que el agente aprenda para optimizar sus decisiones. En relación a esto, se subraya que la dimensión del aprendizaje determina si: el conocimiento es otorgado o si se adquiere a través de datos o experiencias previas.

Las reflexiones de [Russell et al. \(2020\)](#) sobre por qué deseáramos que un agente aprenda en lugar de programarlo correctamente desde el principio son esclarecedoras. Señalan que los diseñadores no siempre pueden prever todas las situaciones futuras y, en muchos casos, ni siquiera saben cómo programar una solución. Por lo tanto, el aprendizaje se convierte en una herramienta esencial para facilitar el diseño del agente.

Los agentes se componen de varios componentes, y estos pueden representarse de diversas maneras dentro del programa del agente. Aunque existe una diversidad en los métodos de aprendizaje, hay una constante unificadora: el aprendizaje en agentes inteligentes puede entenderse como un proceso que modifica cada componente del agente en relación

con la retroalimentación disponible, mejorando, de este modo, el rendimiento global del agente.

2.2.2. Tipos de aprendizaje automático

El vasto dominio del aprendizaje automático se categoriza, en gran medida, por la naturaleza de la retroalimentación y la información disponible. A medida que las máquinas procesan y analizan datos, el tipo de retroalimentación que reciben influye profundamente en el proceso y los resultados del aprendizaje. Según [Russell et al. \(2020\)](#), el espectro del aprendizaje automático se puede clasificar en tres enfoques esenciales basados en los diferentes tipos de retroalimentación que pueden acompañar a los datos de entrada: aprendizaje supervisado, aprendizaje no supervisado y aprendizaje por refuerzo.

Cada uno de estos enfoques posee características y aplicaciones únicas, reflejando distintas metodologías y objetivos en el campo del aprendizaje automático. En las siguientes secciones, se desentrañarán las peculiaridades y fundamentos de estos tres tipos de aprendizaje, proporcionando una visión comprensiva de sus alcances, ventajas y desafíos en el panorama actual de la inteligencia artificial.

2.2.2.1. Supervisado

El aprendizaje supervisado se establece como uno de los paradigmas más fundamentales y prevalentes en el ámbito del aprendizaje automático. Como sugiere su nombre, este método implica la “supervisión” del proceso de aprendizaje de la máquina. ¿Cómo se realiza esta supervisión? Proporcionando al algoritmo un conjunto predefinido de ejemplos, más precisamente, una serie de pares de entrada-salida. Estos pares guían a la máquina para comprender la relación entre la entrada, a menudo denominada atributos, y la salida o su etiqueta. Dichos atributos y etiquetas pueden manifestarse tanto en formas cualitativas, como variables categóricas o discretas, como en formatos cuantitativos.

Esta distinción delinea dos tareas predictivas principales: regresión, que se refiere a la predicción de salidas cuantitativas, y clasificación, que gira en torno a la predicción de salidas cualitativas. Desde una perspectiva formal, el aprendizaje supervisado puede concebirse como la búsqueda para descubrir una función h , adecuadamente denominada hipótesis, de un conjunto de datos que se alinee más estrechamente con una función f verdadera

pero desconocida. La aspiración principal es que h no solo imite f sobre los ejemplos proporcionados, sino que también extienda su aplicabilidad sobre nuevos ejemplos no vistos, asegurando sus capacidades de generalización.

De hecho, el verdadero testimonio de la eficacia de un algoritmo de aprendizaje automático no es simplemente su competencia sobre el conjunto de entrenamiento, sino su habilidad para navegar por territorios inexplorados; datos a los que no ha sido expuesto. Esto requiere la introducción de un conjunto de prueba, que sirve como criterio para medir la precisión de h . Se considera que una hipótesis que puede navegar hábilmente por este conjunto de prueba tiene sólidas habilidades de generalización.

Dentro de esta modalidad de aprendizaje, diversos algoritmos han sido propuestos y ampliamente utilizados en la literatura. A continuación, se describen brevemente algunos de los más destacados:

- **Árboles de Decisión (DT):** Son modelos gráficos que toman decisiones basadas en realizar preguntas sobre los atributos de los datos. Estos árboles dividen el conjunto de datos en subconjuntos más pequeños, facilitando así la tarea de predicción (Quinlan, 1986; Breiman, 2001).
- **Redes Bayesianas (BN):** Estas son representaciones gráficas que muestran las relaciones probabilísticas entre un conjunto de variables (Friedman et al., 1997).
- **Regresiones Lineales (LR):** Son técnicas que buscan encontrar la relación entre una o más variables independientes y una variable dependiente. Se utilizan principalmente cuando la relación entre las variables es aproximadamente lineal (Montgomery et al., 2015).
- **Máquinas de Vectores de Soporte (SVM):** Estos algoritmos buscan encontrar un hiperplano que mejor separe las clases en un espacio de características. Son especialmente útiles en problemas de clasificación (Hearst et al., 1998).
- **Redes Neuronales Artificiales (ANN):** Inspiradas en el funcionamiento del cerebro humano, estas redes son conjuntos interconectados de nodos o “neuronas”. Han demostrado ser altamente efectivas en una variedad de tareas, desde la clasificación hasta la generación de contenido (Haykin et al., 1999).

- **Potenciación del Gradiente (GB):** Es una técnica que construye un modelo predictivo en forma de un conjunto de modelos predictivos débiles, típicamente árboles de decisión (Friedman, 2001).

Sin embargo, por potente que parezca el aprendizaje supervisado, viene atado con su propio conjunto de desafíos y requisitos:

1. Calidad y cantidad de datos: La base del aprendizaje supervisado descansa en el pilar de datos de calidad. La efectividad del algoritmo es directamente proporcional a la calidad y cantidad de ejemplos etiquetados sobre los que se entrena. Datos insuficientes o de baja calidad pueden perjudicar gravemente el rendimiento del modelo.
2. Sobreajuste: Aunque apuntar al mejor ajuste podría parecer el escenario ideal, puede llevar a que el modelo se adapte excesivamente a los datos de entrenamiento, perdiendo así sus habilidades de generalización sobre nuevos datos.
3. Restricciones computacionales: Algunos algoritmos de aprendizaje supervisado, especialmente aquellos más intrincados como las redes neuronales profundas, demandan un considerable poder y recursos computacionales.
4. Interpretabilidad: Mientras que algunos modelos, como las redes bayesianas, ofrecen claridad de forma inherente, otros, en particular modelos complejos como las redes neuronales, pueden ser percibidos como “cajas negras”, lo que los hace más difíciles de interpretar o comprender.

A pesar de estos desafíos, el advenimiento y evolución de técnicas como la regularización, los métodos de conjunto y los avances en hardware solo han reforzado la prominencia del aprendizaje supervisado. Su adaptabilidad a través de una multitud de aplicaciones, desde el reconocimiento de imágenes y voz hasta la predicción financiera, solo consolida su importancia primordial en la evolución de la inteligencia artificial. Como subrayan la gama de algoritmos, el aprendizaje supervisado sigue siendo un pilar en el reino del aprendizaje automático, impulsando innovaciones y soluciones en numerosos dominios.

2.2.2.2. No supervisado

El aprendizaje no supervisado es una rama del aprendizaje automático donde los algoritmos buscan patrones en el conjunto de datos de entrada sin recibir una retroalimentación

explícita. A diferencia del aprendizaje supervisado, no se cuenta con etiquetas o categorías predefinidas para guiar el proceso de aprendizaje. En lugar de ello, se centra en descubrir estructuras ocultas dentro de los datos.

Uno de los objetivos más comunes del aprendizaje no supervisado es la agrupación o *clustering*. Esta tarea consiste en detectar grupos similares dentro de los datos. Formalmente, dada una serie de N observaciones $(x_1, x_2, \dots, x_N)$ de un vector aleatorio p -dimensional X con una densidad conjunta $Pr(X)$, el objetivo es inferir directamente las propiedades (patrones) de esta densidad de probabilidad sin la ayuda de un supervisor que proporcione las respuestas correctas o el grado de error para cada observación.

Mientras que en el aprendizaje supervisado existe una medida clara de éxito que se puede utilizar para juzgar la idoneidad en situaciones particulares y para comparar la eficacia de diferentes métodos, en el contexto del aprendizaje no supervisado no existe tal medida directa de éxito. La validez de las inferencias obtenidas de la mayoría de los algoritmos de aprendizaje no supervisado es difícil de determinar. Se deben emplear argumentos heurísticos no sólo para motivar los algoritmos, sino también para juzgar la calidad de los resultados. Esta situación ha llevado a una gran proliferación de métodos, ya que su eficacia es subjetiva y no puede ser verificada directamente.

Dentro del aprendizaje no supervisado, existen diversos algoritmos que se han consolidado como herramientas esenciales debido a su capacidad para identificar patrones y estructuras ocultas en grandes conjuntos de datos. Aunque cada algoritmo tiene su propia especialidad y ventajas, también vienen acompañados de ciertos desafíos inherentes a sus mecanismos y principios operativos. Entre los algoritmos más destacados, encontramos:

■ **Análisis de Componentes Principales (PCA)** (Abdi y Williams, 2010):

- *Ventajas*: Es un método estadístico que permite simplificar la complejidad en conjuntos de datos multivariados, manteniendo toda la información posible. Es particularmente útil para la reducción de dimensiones.
- *Desafíos*: Interpretar los componentes principales en contextos donde las variables tienen significados complejos o ambiguos puede ser complicado. Además, determinar cuántos componentes principales conservar también puede ser una tarea desafiante.

■ **Algoritmo K-means** (Sculley, 2010; Navarro et al., 1997):

- *Ventajas*: Es un algoritmo de clustering rápido y eficiente, especialmente útil para dividir un conjunto de datos en grupos homogéneos.
- *Desafíos*: La elección del número adecuado de agrupaciones (valor de k) es fundamental para el buen funcionamiento del algoritmo, pero determinar este valor no siempre es trivial. Además, el algoritmo es sensible a la inicialización y puede converger a mínimos locales.

Cada uno de estos algoritmos, con sus ventajas y desafíos, ofrece herramientas únicas para abordar problemas específicos en el ámbito del aprendizaje automático no supervisado. La elección adecuada del algoritmo a utilizar dependerá del tipo de datos con los que se trabaje y del objetivo específico que se quiera alcanzar en el análisis.

Es esencial comprender que el aprendizaje no supervisado no pretende reemplazar a su contraparte supervisada, sino complementarla. Ambos enfoques ofrecen ventajas únicas según la naturaleza del problema y la disponibilidad de datos etiquetados.

2.2.2.3. Por refuerzo

A diferencia de los anteriores enfoques de aprendizaje, en el aprendizaje por refuerzo no se aprende de una base de datos recibida de forma pasiva, ya que en algunas aplicaciones reales, la cantidad de ejemplos necesarios para aprender simplemente no existe (Russell et al., 2020) y esta debe ser recolectada activamente por el aprendiz. Según este enfoque, el algoritmo de Machine Learning adquiere conocimientos a partir de la interacción con el entorno (Sutton y Barto, 2018). El aprendizaje por refuerzo puede ser conceptualizado como el estudio de la planificación y el aprendizaje en un contexto donde un aprendiz interactúa activamente con el ambiente para alcanzar un objetivo definido. Esta interacción activa justifica el uso del término “agente” para referirse al aprendiz. La consecución del objetivo del agente se mide típicamente por la recompensa que recibe del entorno y que este busca maximizar.

El escenario general para el aprendizaje por refuerzo es ilustrado en el Figura 2.4. El agente recopila información a través de una secuencia de acciones al interactuar con el medio ambiente. En respuesta a una acción, el aprendiz o agente recibe dos tipos de información:

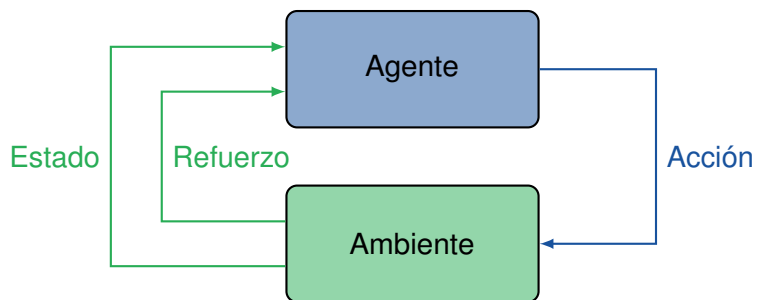


Figura. 2.4: Representación del escenario general del aprendizaje por refuerzo (Sutton y Barto, 2018).

su estado actual en el entorno y una recompensa de valor real, que es específica para la tarea y su objetivo correspondiente.

El aprendizaje por refuerzo se realiza utilizando Procesos de Decisión de Markov (MDPs por sus siglas en inglés). Un MDP se define por un conjunto de estados S , un conjunto de acciones A , una función de transición de estados T , y una función de recompensa R . En cada paso temporal, el agente realiza una acción a_t en el estado s_t , lo que produce una recompensa r_t y un nuevo estado s_{t+1} de acuerdo con la función de transición de estados. El objetivo es aprender una política π que maximice la recompensa a largo plazo esperada. Sin embargo, el agente solo recibe recompensas inmediatas del entorno y no recibe retroalimentación sobre recompensas a futuro o a largo plazo. Por ello, es crucial considerar recompensas o penalizaciones retardadas en el aprendizaje por refuerzo. Esto crea una dicotomía entre explorar estados y acciones desconocidos para obtener más información sobre el entorno y las recompensas, y explotar la información ya adquirida para optimizar la recompensa. Esta dicotomía se conoce como el dilema exploración versus explotación, y es inherente al aprendizaje por refuerzo.

Algunos de los algoritmos de aprendizaje por refuerzo más conocidos son:

■ **Q-learning** (Watkins, 1989):

- Está orientado a la toma de decisiones.
- Aprende la función de valor de acción que le da la máxima recompensa acumulada en el futuro, tomando una acción en un estado particular.
- No requiere un modelo del entorno, y puede manejar problemas con recompensas estocásticas y transiciones no deterministas.

■ **SARSA** (Rummery y Niranjan, 1994):

- SARSA representa “State-Action-Reward-State-Action”.
- Se utiliza para aprender la función de valor de un agente.
- A diferencia del Q-learning, SARSA toma en cuenta la acción actual y la siguiente acción al actualizar el valor Q, lo que lo convierte en un algoritmo basado en políticas “on-policy”, mientras que Q-learning es “off-policy”.

Aunque el aprendizaje por refuerzo es una de las técnicas más prometedoras para revolucionar la Inteligencia Artificial (Russell et al., 2020), lo cierto es que los avances más conocidos en este terreno se han generado en problemas con ambientes con nula relación a los ambientes encontrados en fenómenos reales. La mayoría de esos avances se limitan a juegos, como el de damas (Samuel, 1959), atari (Mnih et al., 2013) y recientemente Go (Silver et al., 2016). Si bien son problemas desafiantes para sus respectivas épocas, esta área se encuentra lentamente transitando a problemas más realistas y más complejos, como en ambientes secuenciales (Leibo et al., 2017), dinámicos y no deterministas (Bloembergen et al., 2015) o parcialmente observables (Omidshafiei et al., 2017; Srinivasan et al., 2018).

Dadas la complejidad y especificidad de algunas simulaciones basadas en agentes, particularmente en las que el entorno y la interacción entre agentes presentan detalles profundos e intrincados, se ha decidido no abordarlas en esta investigación. Se reconoce que, en contextos donde existen múltiples agentes no cooperativos y ambientes parcialmente observables, la convergencia del algoritmo hacia una solución óptima podría nunca materializarse. En consecuencia, la inclusión de tales escenarios podría representar una barrera insuperable para alcanzar resultados concluyentes en este trabajo.

Por lo tanto, en esta investigación, se sugiere emplear técnicas de aprendizaje por refuerzo en el modelado y simulación basado en agentes de manera esquemática. Esto establece directrices sobre las circunstancias en las cuales sería apropiado recurrir al aprendizaje por refuerzo. Sin embargo, queda en manos de investigadores con intereses más específicos, realizar un análisis exhaustivo y riguroso del entorno y las particularidades de su caso de aplicación para elegir el algoritmo de aprendizaje más adecuado.

CAPÍTULO 3

REVISIÓN DE LA LITERATURA

Contenido del capítulo

3.1. Metodología	56
3.1.1. Objetivo de la revisión	56
3.1.2. Organización de la revisión: temas principales	57
3.1.3. Tipo de revisión de la literatura	58
3.1.3.1. Revisión sistemática	58
3.1.3.2. Revisión narrativa tradicional	59
3.1.4. Búsqueda de literatura	59
3.1.5. Criterios de inclusión o exclusión	60
3.2. Revisión de la literatura	62
3.2.1. Análisis bibliométrico	62
3.2.1.1. Análisis estadístico inicial y bibliométrico	63
3.2.1.2. Análisis de estructura conceptual	66
3.2.2. Revisión narrativa	73
3.2.2.1. Aprendizaje <i>offline</i> a priori	76
3.2.2.2. Aprendizaje <i>offline</i> a posteriori	77
3.2.2.3. Aprendizaje <i>online</i>	80
3.3. Discusión	81
3.3.1. Análisis de tendencias	82
3.3.2. Oportunidades emergentes	83
3.3.3. Estándares de documentación	85
3.4. Conclusión	86
3.5. Apéndice	89

El campo interdisciplinario del aprendizaje automático y la modelación y simulación basados en agentes ha experimentado un crecimiento notable en las últimas décadas. Esta convergencia ha suscitado una variedad de investigaciones, propuestas y desarrollos en

el ámbito académico y profesional. Es crucial, por lo tanto, realizar una revisión exhaustiva y crítica de estos avances para identificar tendencias, lagunas, y oportunidades en la investigación.

La revisión de literatura desempeña un papel crucial al proporcionar una comprensión sistemática y contextual de cómo las técnicas de aprendizaje automático se han integrado con la modelación y simulación basadas en agentes. Esta comprensión no solo ofrece una perspectiva histórica sino que, más importante aún, destaca las lagunas, oportunidades y direcciones futuras para la investigación. Es por ello que este capítulo, al intentar identificar y evaluar críticamente los desarrollos existentes en el ámbito, se convierte en una herramienta esencial para los investigadores que buscan construir sobre fundamentos sólidos.

Con esta finalidad, el presente capítulo ofrece una revisión de literatura meticulosa. Aunque se parte de la revisión narrativa tradicional, se han incorporado elementos de la revisión sistemática para garantizar una cobertura exhaustiva y objetiva. El capítulo se estructura en tres secciones distintas: la primera detalla la metodología empleada para guiar la búsqueda, asegurando así la replicabilidad y validez de nuestro enfoque; la segunda sección, que constituye el cuerpo principal de la revisión, desgrana los hallazgos y avances más significativos en el área; y la última sección, después de ofrecer un resumen comprensivo, discute el estado actual del campo, resaltando potenciales áreas de interés para futuras investigaciones.

3.1. Metodología

Para volver transparente el proceso de revisión, en este primer apartado se definen los criterios metodológicos empleados para desarrollar la revisión literaria, involucrando aspectos como el objetivo de la revisión, los temas principales que serán analizados, el tipo de revisión que será empleado, los criterios para incluir o excluir material bibliográfico, así como las palabras clave, bases de datos y parámetros utilizados en la búsqueda de documentos.

3.1.1. Objetivo de la revisión

El objetivo principal de la revisión de la literatura es presentar una imagen integrada del estado actual del conocimiento sobre la incorporación del aprendizaje automático en la modelación y simulación basada en agentes, haciendo hincapié en: 1) los tipos de aprendizaje automático utilizados y su aplicación en diferentes etapas específicas del proceso de

modelación, a decir, diseño, calibración, exploración-experimentación y/o análisis-validación; y 2) la metodología empleada para realizar la integración de ambas herramientas.

Este objetivo se logrará sintetizando la literatura pasada y actual, destacando diferentes enfoques metodológicos, así como estudios relevantes. Por lo que el centro de atención en la revisión será por un lado las prácticas y aplicaciones de la investigación, centrándose en cómo se puede aplicar una determinada técnica de aprendizaje en la modelación y simulación basada en agentes; y por otro, los métodos empleados, destacando las estrategias de documentación, diseño del modelo, calibración, validación, análisis e interpretación de los datos resultantes, así como las fortalezas y debilidades de la metodología elegida.

Esto permitirá identificar los principales problemas, tendencias, complejidades y controversias que están en el centro del tema de estudio, así como una posible dirección para futuras investigaciones, problemas que deben explorarse y atenderse o posibles aplicaciones para la práctica.

3.1.2. Organización de la revisión: temas principales

Como se puede intuir, los dos temas principales bajo los cuales se conduce la revisión de la literatura son aprendizaje automático y modelación y simulación basada en agentes. Sin embargo, se hace una distinción importante sobre el momento en que es aplicado el aprendizaje en el proceso de modelación y simulación, reconociendo principalmente dos momentos, durante la simulación o fuera de ella.

El primer caso, se ha denominado aprendizaje en línea, y aunque puede emplearse cualquier tipo de aprendizaje automático, tiene una correspondencia más natural con el aprendizaje por refuerzo. El aprendizaje en esta modalidad tiene como propósito la exploración dentro de la modelación y simulación basada en agentes, ya que se espera una política “óptima” de desempeño por parte del o de los agentes así como conocimiento del ambiente.

En lo que respecta al segundo caso, llamado aprendizaje fuera de línea, este puede dividirse en etapas diferentes dentro del modelado y simulación basado en agentes, en primer lugar hacia aquel aprendizaje destinado al diseño del mismo modelo o a la calibración, que se denomina aprendizaje “a priori”, y en segundo lugar, aquel aprendizaje que se

implementa para realizar la validación del modelo o el análisis de los resultados, en cuyo caso es nombrado aprendizaje “a posteriori”.

3.1.3. Tipo de revisión de la literatura

Según Efron y Ravid (2018), los tipos de revisiones de la literatura se presentan en un continuo que va por un lado desde un marco científico-cuantitativo representado por revisiones sistemáticas, hasta un marco interpretativo-cualitativo, ejemplificado en una revisión hermenéutica-fenomenológica, del otro lado del continuo. Entre estos dos tipos opuestos de revisiones de la literatura se encuentra la revisión narrativa tradicional que integra diferentes enfoques de investigación. A continuación se describen brevemente las características de la revisión sistemática y la revisión narrativa tradicional.

3.1.3.1. Revisión sistemática

En cuanto a la revisión sistemática, es un enfoque que está altamente estructurado y basado en protocolos. Los defensores de este estilo de revisión afirman que es imparcial, sistemático, riguroso y replicable. El propósito de la revisión sistemática es responder una pregunta bien enfocada y específica que se formula antes de emprender la búsqueda en la biblioteca. El investigador realiza una búsqueda rigurosa y exhaustiva con el fin de identificar todos los estudios posibles y relevantes sobre el tema que investiga.

La búsqueda se basa en un protocolo estricto y explícito que evalúa y sintetiza la evidencia empírica reportada en estudios individuales. Se requiere que el autor de la revisión sistemática sea neutral y objetivo para minimizar sesgos y errores. Se formulan criterios de exclusión e inclusión predeterminados para garantizar que la información obtenida de las fuentes sea precisa e imparcial y que se utilicen métodos bien definidos para evaluar los resultados de cada estudio.

En sus escritos, los investigadores de revisiones sistemáticas evalúan y sintetizan la evidencia empírica reportada en estudios individuales para llegar a respuestas concluyentes a la pregunta de investigación. Aunque una revisión sistemática puede incluir alguna investigación cualitativa, la mayoría de las revisiones de este tipo son cuantitativas y utilizan datos estadísticos.

3.1.3.2. Revisión narrativa tradicional

El estilo de revisión narrativa tradicional sigue siendo el método más común entre estudiantes e investigadores en ciencias sociales y educación. La revisión narrativa tradicional se basa en una variedad de disciplinas académicas e incluye diversos métodos de investigación, estudios cualitativos, cuantitativos y teóricos. Este tipo de revisión examina el estado del conocimiento en un área temática específica y ofrece una base integral para comprender ese tema en particular.

Este estilo resume de manera crítica las teorías, examina los estudios e investiga los métodos utilizados en la investigación existente. El revisor recopila un amplio espectro de la literatura escrita sobre el tema y lo sintetiza en una interpretación coherente que resalta los principales problemas, tendencias, complejidades y controversias que están en el centro del mismo. El autor también puede identificar una posible dirección para futuras investigaciones, problemas que deben explorarse o posibles aplicaciones para la práctica.

Tipo de revisión de la literatura elegido

Para los objetivos de la presente investigación, se adopta un tipo de revisión dentro del continuo, que parte de la revisión narrativa tradicional y que incorpora algunos elementos de la revisión sistemática, dando a conocer la fórmula de búsqueda del material bibliográfico, así como los criterios de inclusión y exclusión, con lo cual se pretende realizar una revisión transparente y replicable.

3.1.4. Búsqueda de literatura

Las Figuras 3.1 y 3.2 muestran las fórmulas empleadas para la búsqueda de literatura en las bases de datos bibliográficos Scopus y Web of Science Core Collection. Las cadenas de texto utilizadas como criterios de búsqueda que refieren a modelo o simulación son buscadas en el título de las publicaciones, mientras que aquellas que refieren a las técnicas de aprendizaje automático son buscadas ya sea en el título, en el resumen o las palabras clave.

La búsqueda literaria arrojó 340 registros en Scopus y 170 en Web of Science. Estos resultados se exportaron en formato CSV (valores separados por comas) para recopilar información esencial de los artículos, como el título, nombres y afiliaciones de los autores,

TITLE (("agent-based model*" OR "individual-based model") OR ("agent-based simulat*" OR "individual-based simulat"))

AND TITLE-ABS-KEY ("machine learning" OR "neural network*" OR {NN} OR {RNN} OR {ANN} OR "decision tree" OR "random forest" OR "logistic regression" OR "Gradient Boosting" OR "Support Vector Machines" OR {SVM} OR "supervised learning" OR "Nearest Neighbor" OR "Naive Bayes" OR "bayesian network" OR "unsupervised learning" OR "k-mean" OR "k-median" OR "Association Rules" OR {RL} OR {MARL} OR "Temporal Difference" OR "Q-learning" OR {SARSA})

Figura. 3.1: Fórmula Scopus.

TI=(("agent-based model*" OR "individual-based model*") OR ("agent-based simulat*" OR "individual-based simulat*"))

AND (TS=("machine learning" OR "neural network*" OR {NN} OR {RNN} OR {ANN} OR "decision tree" OR "random forest" OR "logistic regression" OR "Gradient Boosting" OR "Support Vector Machines" OR {SVM} OR "supervised learning" OR "Nearest Neighbor" OR "Naive Bayes" OR "bayesian network" OR "unsupervised learning" OR "k-mean" OR "k-median" OR "Association Rules" OR {RL} OR {MARL} OR "Temporal Difference" OR "Q-learning" OR {SARSA}))

Figura. 3.2: Fórmula Web of Science Core Collection.

resumen, palabras clave y referencias. Estos archivos se emplearán en análisis posteriores. El primer paso es identificar y eliminar registros duplicados. Tras esta depuración, en la que se eliminaron 130 registros, se contabilizaron 380 documentos en esta fase inicial.

3.1.5. Criterios de inclusión o exclusión

Para asegurar la relevancia y calidad de los artículos seleccionados en este estudio, se aplicaron diversos criterios de inclusión y exclusión en la revisión de la literatura:

C1 Especificidad y enfoque de investigación: Es vital incluir solo publicaciones que presenten propuestas concretas vinculadas a problemas de ABM abordados con una perspectiva de ML. Esta elección asegura que los resultados sean directamente aplicables y pertinente a la investigación principal, y que las publicaciones tengan un grado de experimentación asociado que respalde las propuestas.

C2 Idioma y alcance de publicación: Se dieron prioridad a las publicaciones internacionales

escritas en inglés, ya que este es el idioma predominante en la literatura científica y garantiza una mayor audiencia y rigurosidad en los procesos de revisión por pares.

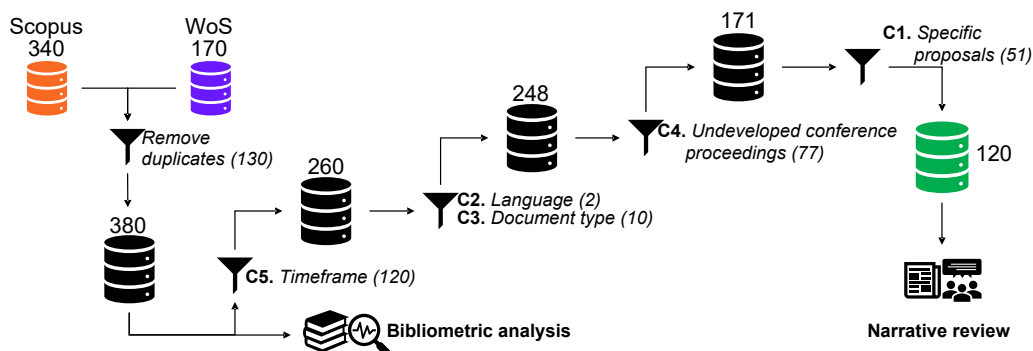
C3 Tipo de publicación: Se optó por excluir capítulos de libros y editoriales para centrarse en fuentes que generalmente tienen un proceso de revisión más riguroso y son más susceptibles de contener investigaciones originales y detalladas.

C4 Naturaleza de la fuente: Las actas de conferencias suelen presentar trabajos en etapas preliminares. Por lo tanto, no se consideraron en la búsqueda inicial, pero se permitió su inclusión en búsquedas complementarias si mostraban un desarrollo y aporte significativos al campo.

C5 Actualidad de la investigación: El campo de la IA está en constante evolución, por lo que se consideró esencial limitar las publicaciones al intervalo de 2015-2022 para asegurar la relevancia y actualidad de los trabajos revisados.

La aplicación del criterio C5 resultó en la eliminación de 120 artículos publicados antes de 2015. Seguidamente, al implementar los criterios C2 y C3, se descartaron 12 artículos. A partir del criterio C4, se excluyeron inicialmente 85 trabajos, pero 8 fueron reconsiderados por su relevancia, incluyendo una revisión de literatura (Ashreeta et al., 2019). Finalmente, tras una evaluación detallada de los 171 artículos que quedaban según el criterio C1, se seleccionaron 120 para ser incluidos en la base de datos. La Figura 3.3 ilustra de forma esquemática este proceso de selección, revisión y análisis de la literatura.

Figura. 3.3: Flujo de revisión de la literatura.



Fuente: Platas-López et al. (2023b).

3.2. Revisión de la literatura

En esta sección se presenta el cuerpo principal de la revisión de la literatura, siguiendo los criterios metodológicos establecidos en la sección anterior. En primer lugar se presenta un análisis bibliométrico con el objetivo de reseñar a grandes rasgos algunas características del cuerpo bibliográfico, lo que nos permitirá agrupar la literatura y posteriormente refinar los criterios de inclusión para proceder con la revisión a profundidad. En segundo lugar, se presenta la revisión de la literatura, el cual deriva del flujo de revisión adoptado para su análisis (Figura 3.3).

3.2.1. Análisis bibliométrico

El análisis bibliométrico es definido como una evaluación cuantitativa de artículos científicos, libros o capítulos de libros publicados (Pritchard et al., 1969). A diferencia de las revisiones sistemáticas de literatura, una revisión bibliométrica tiene la facilidad de proporcionar información sobre dominios caracterizados por grandes¹ cantidades de información bibliográfica (Goodell et al., 2021). Tomando en consideración esta ventaja, el análisis bibliométrico tiene el potencial para refinar los criterios de inclusión/exclusión sobre los registros encontrados.

Existen diferentes alternativas de software para el análisis bibliométrico, cada uno con diferentes capacidades y limitaciones. Algunas de las herramientas más populares incluyen Publish or Perish², HistCite (Garfield et al., 2006), BibExcel (Persson, 2008) y un paquete de R (R Core Team, 2021) denominado bibliometrix (Aria y Cuccurullo, 2017).

Se eligió bibliometrix para este estudio debido a que al operar en el entorno de R ofrece un alto grado de flexibilidad para modificar y/o ajustar los datos de entrada importados de varias bases de datos, incluidas Scopus y Web of Science, y otras como Dimensions, Lens.org, PubMed o Cochrane library, así como la capacidad de proporcionar un análisis de datos completo para su uso en una variedad de herramientas de análisis de red, sin la necesidad

¹ El término “grandes” es relativo y puede variar dependiendo de los recursos disponibles. Si se cuenta con más tiempo y personal, se puede abordar un número mayor de documentos. Sin embargo, con recursos limitados, el análisis de la misma cantidad de documentos podría ser prohibitivo. Por ello, es fundamental establecer criterios claros para delimitar cuánto se puede procesar en función de los recursos disponibles.

² <https://harzing.com/resources/publish-or-perish>

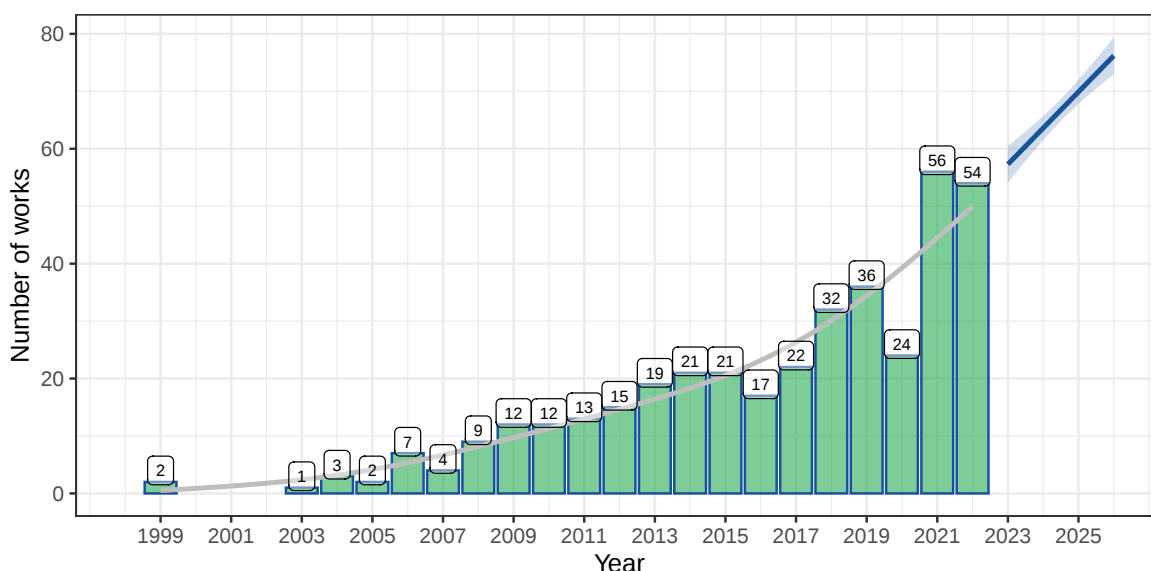
de exportar a otras herramientas como Pajek (Batagelj y Mrvar, 2002), Gephi (Bastian et al., 2009) o VOSviewer (van Eck y Waltman, 2010).

Ninguna de las herramientas para análisis bibliométrico mencionadas proporciona datos de análisis de red suficientemente detallados en su salida. Por lo tanto, bibliometrix es una herramienta poderosa para analizar datos bibliográficos, con un entorno operativo relativamente sencillo. Usamos bibliometrix para realizar un análisis bibliométrico y estadístico inicial, así como un análisis de red adicional.

3.2.1.1. Análisis estadístico inicial y bibliométrico

Los registros hallados en la búsqueda fueron publicados entre el año 1999 y 2022. La Figura 3.4 muestra una tasa de crecimiento anual de 15.41 % en la cantidad de material bibliográfico publicado sobre modelación y simulación basada en agentes y aprendizaje automático. De persistir la misma tendencia, se estima con un modelo aditivo generalizado, que para el año 2026 serán publicados entre 73 y 78 artículos, capítulos o libros sobre este tema.

Figura. 3.4: Tendencia de publicación sobre modelación y simulación basada en agentes y aprendizaje automático.



Fuente: Elaboración propia.

Las estadísticas iniciales muestran que 205 fuentes han contribuido a la publicación los 380 registros bibliográficos. Se encontró que 10 de estas fuentes han publicado 73 de

estos artículos identificados, lo que representa el 19.21 % de todos los registros publicados, aproximadamente la quinta parte. La Tabla 3.1 muestra las fuentes en las que aparecieron estas publicaciones.

Tabla 3.1: Las 10 principales fuentes editoriales que contribuyen al área de modelación y simulación basada en agentes y aprendizaje automático.

Fuentes	Publicaciones
Lecture Notes in Computer Science	21
Ecological Modelling	17
Plos ONE	6
Physica A: Statistical Mechanics and its Applications	5
Communications in Computer and Information Science	4
Environmental Modelling and Software	4
International Conference on the European Energy Market	4
Journal of Theoretical Biology	4
Proceedings - Winter Simulation Conference	4
Proceedings of the International Joint Conference AAMAS	4

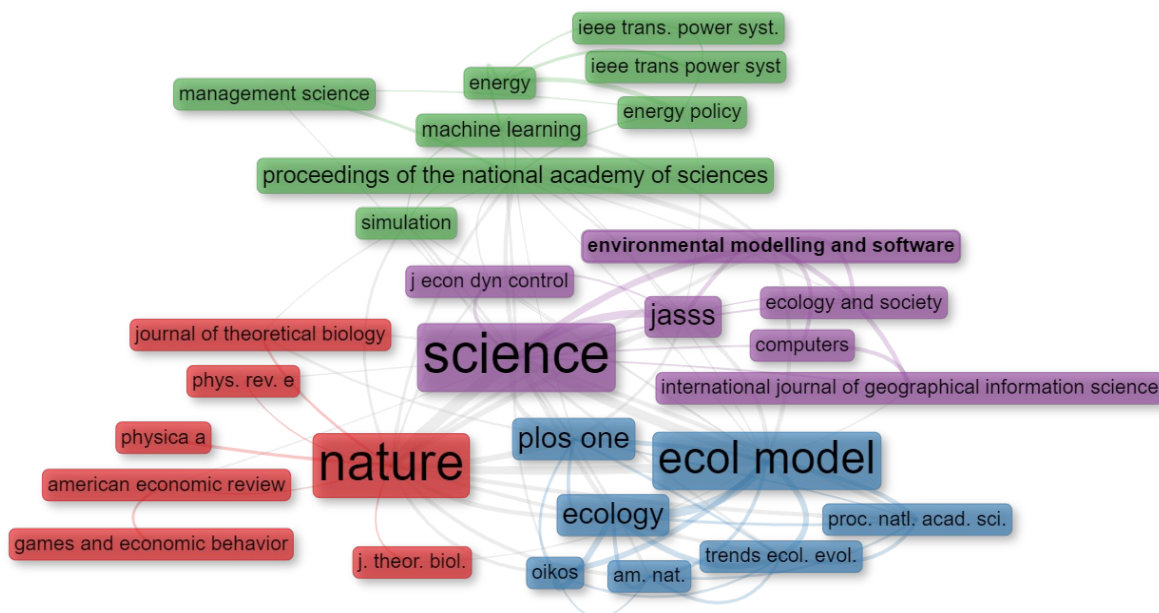
Además, los datos sobre citas entre revistas muestran que Ecological Modelling, Journal of Artificial Societies and Social Simulation, Science y Nature son las más citadas entre los documentos encontrados. La Fig. 3.5 muestra la relación y la magnitud de dicha relación entre las revistas más comunes, tomando como indicador el número de referencias que aparecen en cada documento.

El campo del autor se extrajo del archivo de datos y se registró la frecuencia de aparición de todos los autores. La Figura 3.6 describe los veinte principales autores que contribuyeron en el área y la cantidad de artículos en que fueron autores o coautores así como el promedio de citas por año. Como puede verse en estos datos, An G., Hayashida T. y Zhang H. parecen dominar la lista.

Para identificar la producción científica a nivel de países, del campo afiliación se extrae el país al que pertenecen los autores y se contabiliza su frecuencia de aparición. El mapa de la Figura 3.7 muestra la nacionalidad de los autores que contribuyeron en el área y la cantidad de documentos en los que participaron. Estados Unidos (231), China (131), Reino Unido (88), Canadá (71) y Japón (47) son los países que dominan.

Cabe destacar que en cuanto a la nacionalidad de los autores correspondientes, nuevamente Estados Unidos (40), China (27), Canadá (23), Reino Unido (12) y Japón (10), son los países que lideran la lista. Se hace una distinción entre los documentos publicados

Figura. 3.5: Red de citación conjunta entre revistas.



Fuente: Elaboración propia.

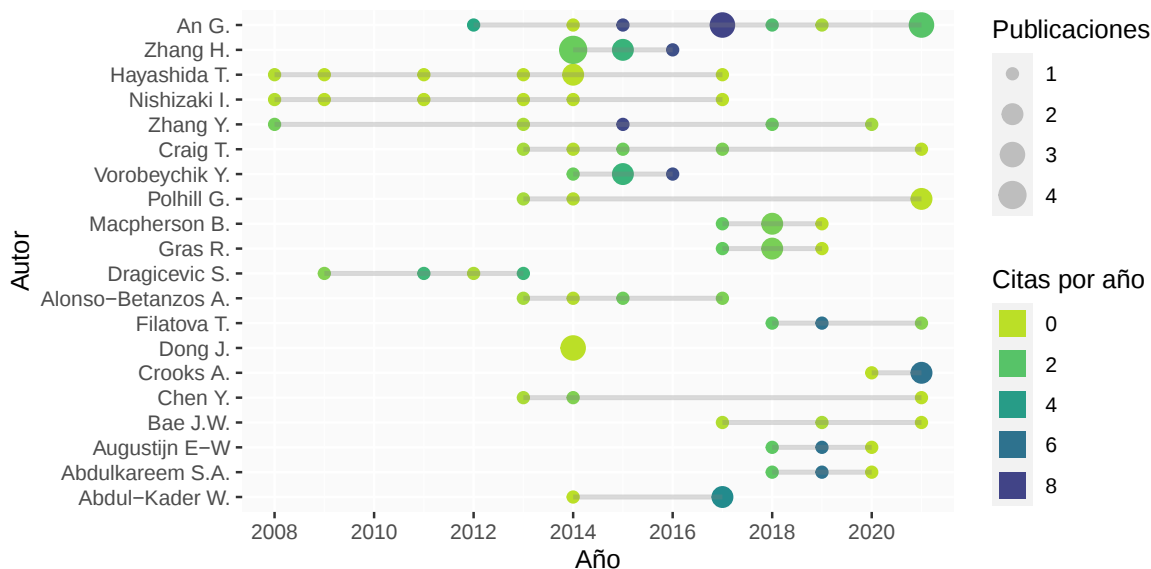
por autores de una sola nacionalidad (PSN), así como aquellas publicaciones en los que intervienen autores de varias nacionalidades (PVN). La Tabla 3.2 muestra a los 10 países que lideran la producción como autores correspondientes, así como la cantidad de publicaciones que se da en colaboración con otros países y de manera independiente.

Tabla 3.2: Número de documentos publicados por país del autor correspondiente. PSN: Publicaciones de una sola nacionalidad. PVN: Publicaciones de varias nacionalidades.

País	Documentos	PSN	PVN	Relación VP
Estados Unidos	40	35	5	0.125
China	27	21	6	0.222
Canadá	23	16	7	0.304
Reino Unido	12	10	2	0.167
Japón	10	8	2	0.200
Alemania	6	4	2	0.333
Irán	6	3	3	0.500
Italia	6	5	1	0.167
España	6	4	2	0.333
Francia	5	2	3	0.600

Para identificar los flujos de colaboración entre países, nuevamente se obtiene la nacionalidad de los autores correspondientes y se hace un nexo con la nacionalidad de los co-autores. La Fig. 3.8 resume esta información mediante un diagrama circular de flujo. Con la finalidad

Figura. 3.6: Top 20 de principales autores contribuyentes, número de artículos publicados y citas promedio anuales.



Fuente: Elaboración propia.

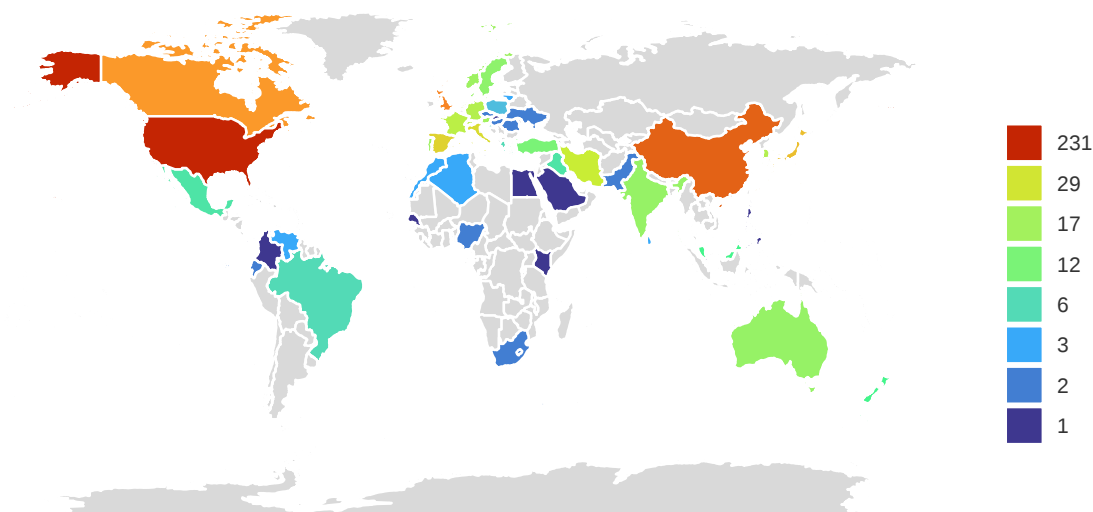
de mejorar la visualización de la información, la nacionalidad de los autores con un sólo artículo en co-autoría, fue concentrada con el nombre de país “Otros”. Aquí se agrupan Brasil, Colombia, Grecia, Hong Kong, Kenia, Malasia, Panamá, Arabia Saudita, Singapur, Eslovenia y Venezuela.

Se puede apreciar que Estados Unidos ha tomado el rol de autor correspondiente en todos los documentos en los que ha trabajado en colaboración con otros países, ya que todos los flujos salen de Estados Unidos pero ninguno entra a ese país. En diferente magnitud y proporción, Reino Unido y China también han asumido ese rol en los documentos en colaboración en los que han participado. También se puede apreciar que China sólo colabora con Estados Unidos.

3.2.1.2. Análisis de estructura conceptual

Para delinear la estructura conceptual del tema, se empleará un mapa temático, en el análisis bibliométrico, un mapa temático representa las relaciones entre términos clave en un conjunto de documentos, utilizando diferentes colores para mostrar la frecuencia o co-ocurrencia de los términos. Consiste en un análisis de redes de co-ocurrencia de palabras para definir sobre qué está discutiendo el campo de investigación, los temas principales

Figura. 3.7: Cantidad de documentos sobre modelación y simulación basada en agentes y aprendizaje automático publicados por país.



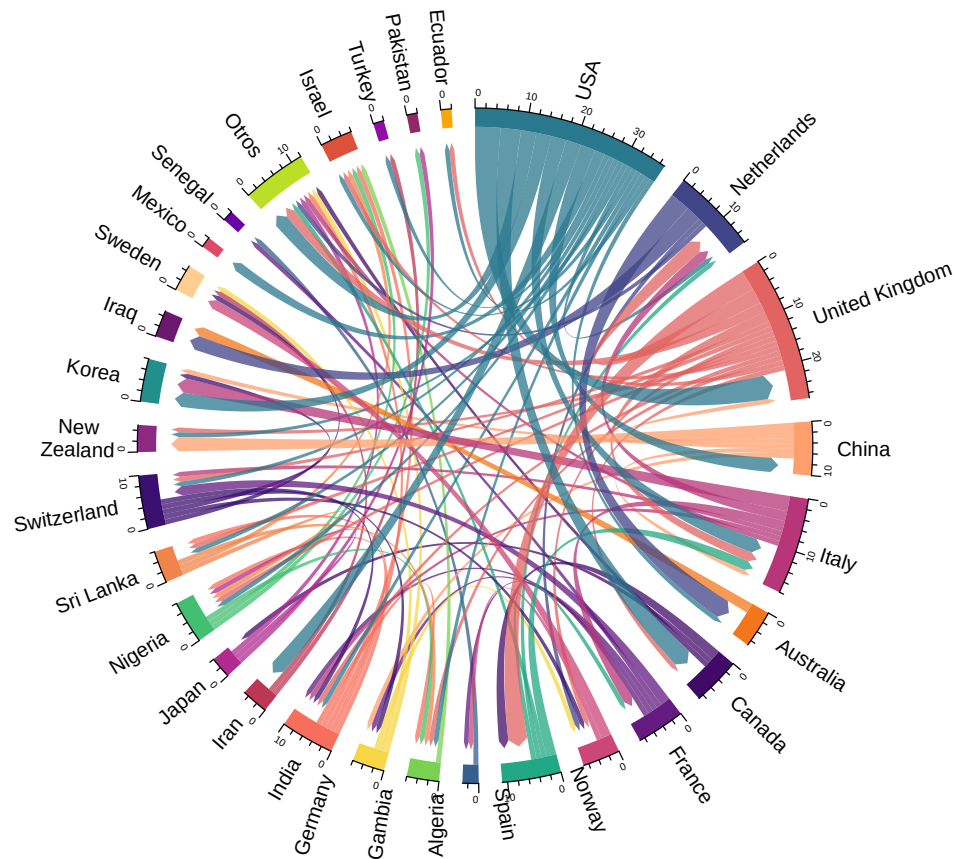
y las tendencias. Interpretar un mapa temático basado en densidad y centralidad implica identificar los términos más relevantes o importantes en el conjunto de documentos, así como las relaciones entre ellos. La densidad se refiere a la frecuencia con la que aparecen términos clave en los documentos, mientras que la centralidad se refiere a la posición relativa de los términos en la red de relaciones entre ellos. Los términos más centrales son aquellos que están más conectados con otros términos, sugiriendo una mayor relevancia o importancia en el conjunto de documentos. Existen cuatro tipos de temas que se refieren a diferentes categorías de tópicos o áreas de investigación en un campo específico:

Temas de nicho. Se caracterizan por tener alta densidad y baja centralidad. Esto significa que hay muchos artículos o publicaciones sobre un tema específico, pero no son citados con frecuencia ni están conectados a otros temas en el campo.

Temas motor. Poseen alta densidad y alta centralidad. Esto implica que hay muchos artículos o publicaciones sobre un tema específico, y son citados con frecuencia o están conectados a otros temas en el campo. Los temas motor son a menudo considerados la fuerza impulsora detrás de la investigación en un campo.

Temas emergentes o en declive. Tienen baja densidad y baja centralidad. Esto significa que hay relativamente pocos artículos o publicaciones sobre un tema específico y

Figura. 3.8: Diagrama circular de flujo de colaboración entre países.



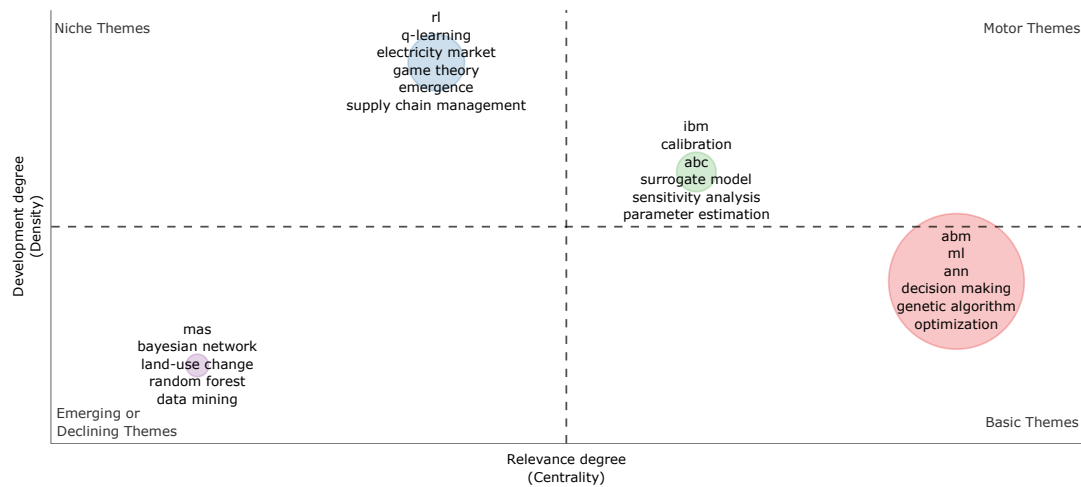
aquellos que existen son citados con poca frecuencia o no están conectados a otros temas en el campo. Los temas emergentes son aquellos que están ganando importancia y atención en el campo, mientras que los temas en declive son aquellos que alguna vez fueron importantes pero ahora están perdiendo relevancia.

Temas base. Poseen baja densidad y alta centralidad. Esto indica que hay relativamente pocos artículos o publicaciones sobre un tema específico, pero son citados con frecuencia o están conectados a otros temas en el campo. Los temas base son a menudo considerados conceptos o teorías fundamentales en un campo.

El mapa temático de la literatura seleccionada sobre ML aplicado a ABM se muestra en la Figura 3.9. Fue calculado utilizando las palabras clave proporcionadas por los autores porque capturan el contenido de los artículos con mayor variedad y precisión que las palabras clave estandarizadas proporcionadas por los editores, las cuales se centran en menos temas relacionados con la computación y simulación, sin diferencias en la estructura del área. Vale

la pena señalar que las palabras clave de los autores fueron lematizadas manualmente para reducirlas a su forma base, estandarizarlas y facilitar su análisis y comparación. Por ejemplo, expresiones comunes o algoritmos se redujeron a su forma abreviada, como ABM, ML o RL para el aprendizaje por refuerzo.

Figura. 3.9: Mapa temático.



Fuente: Platas-López et al. (2023b).

Los temas de nicho en el cuadrante superior izquierdo también son conocidos como temas altamente desarrollados y aislados. Tienen conexiones internas bien desarrolladas (alta densidad) pero casi ninguna conexión externa y, por lo tanto, tienen una importancia limitada para el campo (baja centralidad). El aprendizaje por refuerzo, Q-learning, mercado eléctrico, teoría de juegos, emergencia y gestión de la cadena de suministro son cuestiones aisladas que han sido ampliamente estudiadas en campos específicos. Sin embargo, integrarlos en un marco comprensivo sigue siendo un desafío.

Los temas emergentes o en declive en el cuadrante inferior izquierdo tienen baja centralidad y densidad, lo que significa que están débilmente desarrollados y son marginales o con poca producción. Podemos ver que las referencias a sistemas multi-agente (MAS) y cambio de uso del suelo están en declive. Sin embargo, es importante señalar que la investigación sobre ABM y MAS ha avanzado de manera independiente, incluso si están estrechamente relacionados. Por otro lado, conceptos emergentes como minería de datos, redes bayesianas y algoritmos de bosque aleatorio están ganando cada vez más atención en la literatura.

Los temas básicos en el cuadrante inferior derecho también se conocen como temas centrales o transversales. Se caracterizan por alta centralidad y baja densidad. Estos temas generales son importantes para un campo de investigación y abarcan las diferentes áreas del campo. Basándonos en los términos de búsqueda, es evidente que los temas que destacan en el cuadrante de temas básicos están relacionados con ABM y ML, lo que indica que están altamente conectados con otros temas y aparecen con frecuencia en la literatura. Sin embargo, vale la pena señalar que otros términos como redes neuronales artificiales, algoritmos genéticos, toma de decisiones y optimización también aparecen en el cuadrante, aunque con menor frecuencia. Estos temas pueden no tener tantas conexiones directas con otros temas en la red, lo que resulta en una menor densidad. Sin embargo, siguen siendo relevantes para el tema general y sugieren que hay investigación y exploración en curso en áreas relacionadas, lo que podría llevar a nuevos desarrollos y avances en estos campos.

Los temas motor en el cuadrante superior derecho se caracterizan por su alta centralidad y densidad. Esto significa que están desarrollados e importantes para el campo de investigación. Entre los temas motor más desarrollados en la literatura, el método de computación bayesiana aproximada (del inglés, Approximate Bayesian Computation, ABC), utilizado para la estimación de parámetros, ha sido de particular interés en el área biológica, donde se acuñó el término modelado basado en individuos (del inglés, Individual-Based Modeling, IBM) (Judson, 1994; Grimm, 1999). En general, es posible notar que el interés está relacionado con aplicaciones de aprendizaje offline aplicadas después de la simulación, en tareas de calibración y análisis.

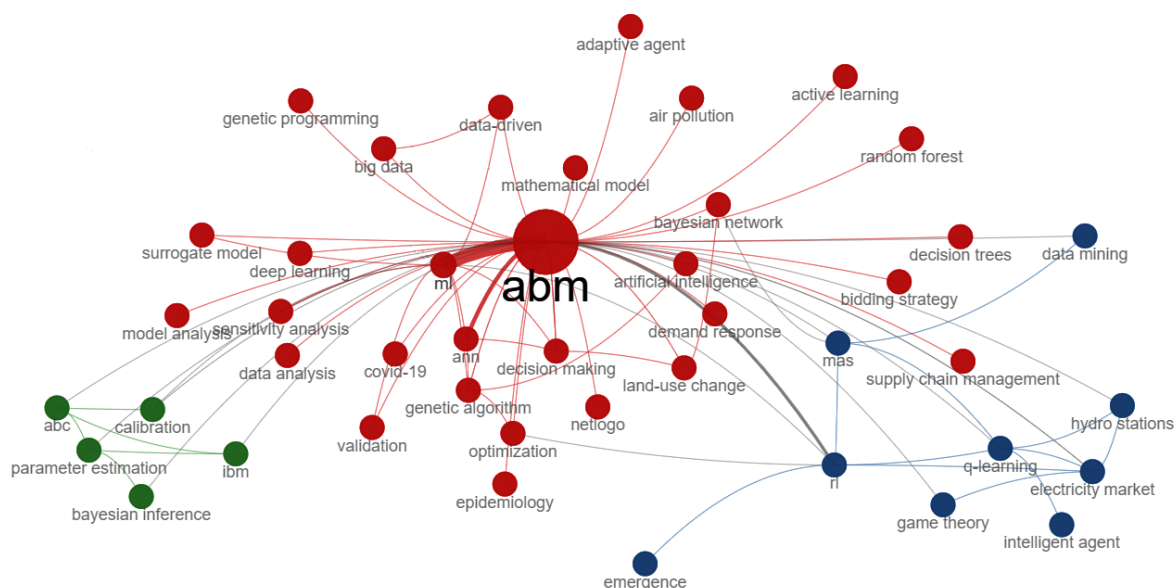
El mapa temático proporciona una visión clara y estructurada de la interacción entre el Aprendizaje Automático y el Modelado Basado en Agentes en la literatura científica. Se observa una variabilidad temática que abarca desde áreas bien establecidas hasta zonas emergentes que todavía requieren una mayor integración en el campo de estudio. Es notable que, si bien hay esfuerzos consolidados en ciertos temas, aún existen oportunidades significativas para integrar y expandir el conocimiento en áreas menos exploradas.

Los esfuerzos actuales en la literatura parecen estar distribuidos, lo que sugiere la necesidad de un enfoque más unificado, especialmente al considerar el potencial sinérgico de combinar ABM y ML. Además, es evidente que, mientras algunos temas han recibido considerable atención, otros, posiblemente igual de relevantes, están emergiendo o decli-

nando, lo cual indica áreas de investigación futura y puntos de enfoque potenciales para académicos y profesionales. En resumen, este análisis cartográfico revela el paisaje actual del interjuego entre ML y ABM, destacando tanto las fortalezas en la investigación existente como las oportunidades para futuras investigaciones.

Ahora, empleando nuevamente las palabras clave de los autores, y para comprender si los temas de ABM y ML interactúan con otros términos y principalmente la forma en que tales términos interactúan, se utiliza un análisis conjunto de palabras a través de la red de palabras clave proporcionadas por los autores. La Figura 3.10 muestra la estructura conceptual entre los problemas más recurrentes abordados por los investigadores en el área. A partir de este análisis, los algoritmos de ML utilizados en el área, como redes neuronales artificiales, computación bayesiana aproximada, algoritmos genéticos, redes bayesianas, bosque aleatorio y Q-learning.

Figura. 3.10: Palabras clave más recurrentes. El ancho de las líneas que unen a los términos es proporcional a su frecuencia de interacción.



Fuente: Platas-López et al. (2023b).

De este análisis, podemos distinguir los algoritmos empleados en el área, destacan por su frecuencia las redes neuronales, el algoritmo de aprendizaje por refuerzo *Q-learning*, arboles de decisión, bosques aleatorios, algoritmos genéticos, redes bayesianas, computación bayesiana aproximada y regresión logística. A su vez, cada uno de estos algoritmos se

encuentra ligado a etapas específicas del proceso de modelación y simulación basado en agentes, o a dominios de aplicación concretos.

Por ejemplo, el algoritmo de computación bayesiana aproximada tiene una relación cercana con la estimación de parámetros o la calibración, y esta última a su vez con el análisis de sensibilidad, lo que denota una preocupación por entender el efecto de los parámetros a su respectiva salida. Y también es posible observar que ésta técnica es más frecuentemente empleada en el área biológica, por su relación con el término modelo basado en individuos.

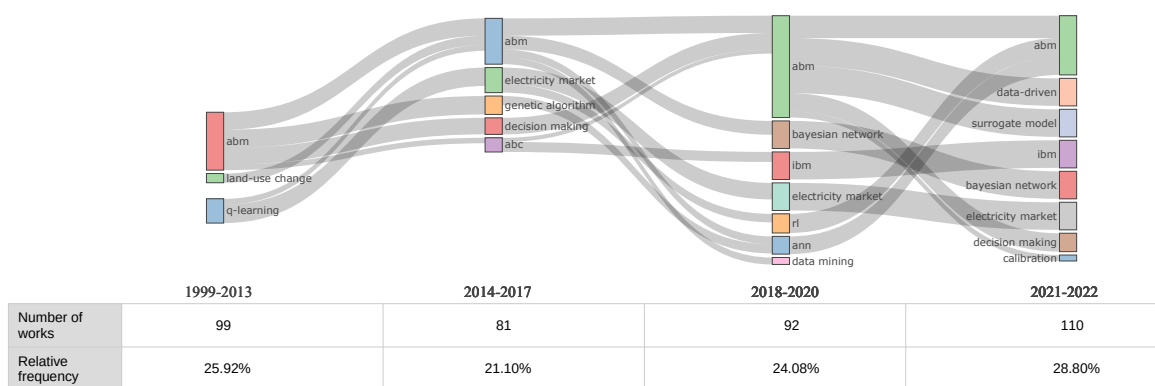
El nexo del aprendizaje por refuerzo se encuentra con algunos dominios específicos de aplicación como lo son el mercado de la electricidad y las estaciones hídricas, mientras que los algoritmos genéticos lo están con problemas de optimización. También es importante destacar que la palabra netlogo se encuentra como uno de los tópicos más empleados por los autores, lo que significa que NetLogo ([Wilensky y Rand, 2015](#)) es un entorno de desarrollo ampliamente usado para la modelación basada en agentes.

Después de comprender la interacción entre palabras clave y los temas dominantes en el área, es esencial examinar cómo estos temas han evolucionado con el tiempo. La dinámica del campo puede revelar tendencias y cambios en el foco de investigación a lo largo de los años, proporcionando una perspectiva más completa de la dirección del área.

Para lograr una visión equilibrada y coherente de esta evolución, se seleccionaron periodos de tiempo que contuvieran una cantidad similar de artículos. Esta elección se basa en el objetivo de asegurar que cada segmento temporal tenga un peso representativo en el análisis, evitando así sesgos que podrían surgir al comparar periodos con volúmenes de publicaciones muy dispares. Por eso, se decide abordar un análisis de evolución temática que segmenta los documentos según cuatro periodos de tiempo determinados, los cuales se representan en el diagrama Sankey ([Schwabish, 2021](#), p. 126) de la Fig. 3.11 en donde también se muestran el número de trabajos publicados y su frecuencia relativa.

Cabe destacar que como era de esperarse, el tema de ABM se mantiene a lo largo del período de tiempo considerado. El resto de los términos pueden dividirse en conceptos básicos derivados de la consulta de búsqueda, por ejemplo, los algoritmos de ML involucrados, y conceptos relacionados con las etapas del ciclo de desarrollo de ABM y sus dominios de aplicación.

Figura. 3.11: Evolución temática.



Fuente: Platas-López et al. (2023b).

En el primer período (1999–2013), los estudios sobre cambios en el uso de la tierra eran populares, y el Q-learning, un método de aprendizaje por refuerzo propuesto en 1989 (Watkins, 1989), era bastante común. Estos términos no aparecen nuevamente en los períodos siguientes. En el segundo período (2014–2017), los estudios sobre mercados eléctricos y toma de decisiones se volvieron populares, mientras que los métodos de ML comúnmente adoptados fueron algoritmos genéticos y la computación bayesiana aproximada (CBA). El dominio de aplicación del mercado eléctrico es recurrente en los períodos siguientes, mientras que la toma de decisiones vuelve a ser relevante en el tercer período (2021–2022).

El dominio del mercado eléctrico está fuertemente relacionado con el aprendizaje por refuerzo, lo cual se manifiesta en la transición al tercer período. La computación bayesiana aproximada es ampliamente utilizada en modelos basados en individuos, como se observa en el tercer y cuarto período. De hecho, no hay dominios de aplicación emergentes en estos períodos, y el interés parece estar centrado en algoritmos de ML, por ejemplo, redes bayesianas y redes neuronales artificiales. En términos de propósitos de aplicación, el enfoque ha sido en métodos basados en datos para diseñar modelos a partir de datos y modelos sustitutos para análisis y calibración.

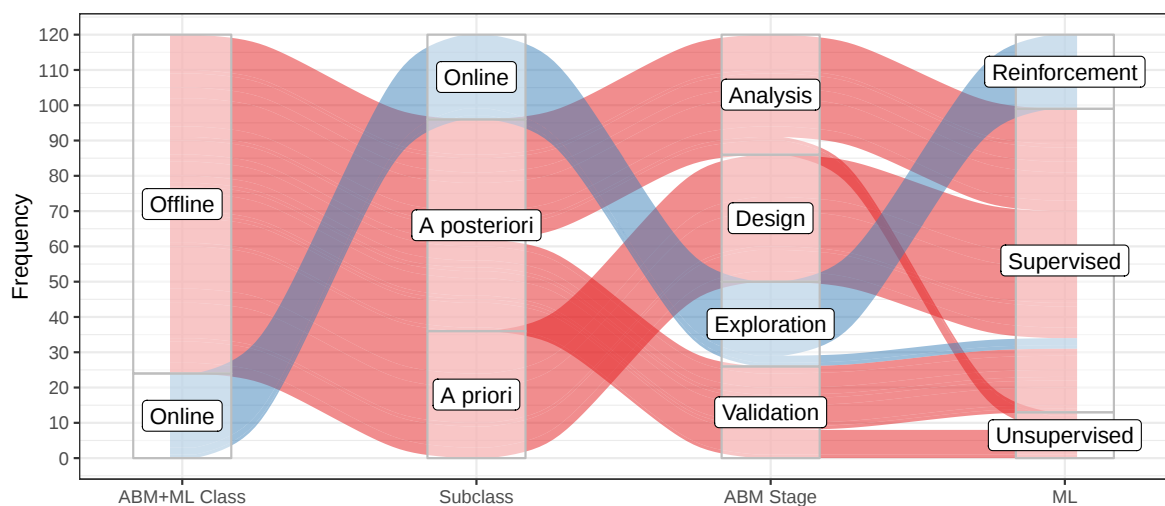
3.2.2. Revisión narrativa

Esta sección recoge los principales resultados de la revisión de 120 soluciones de ML en ABM en los últimos años, es decir, entre 2015 y 2022. Por ello, se proporciona una revisión

actualizada sobre la incorporación de técnicas de ML en las diferentes etapas del proceso de modelado basado en agentes. Para mejorar algunos procesos en ABM, ML ha jugado un papel esencial en el paso de construcción y análisis del modelo.

Los tres principales paradigmas de aprendizaje son el aprendizaje supervisado, el aprendizaje no supervisado y el aprendizaje por refuerzo. La figura 3.12 resume los trabajos revisados por la taxonomía propuesta y el tipo de aprendizaje que emplean. El aprendizaje por refuerzo se usa ampliamente para el aprendizaje en línea, mientras que los otros dos paradigmas de ML se han usado ampliamente para el aprendizaje fuera de línea. Aunque el aprendizaje no supervisado tiene solo dos aplicaciones.

Figura. 3.12: Artículos revisados sobre ML en ABM por etapa de aplicación y tipo de algoritmo de ML.



Fuente: Platas-López et al. (2023b).

Complementario a la representación gráfica, la Tabla 3.3 enumera las 120 publicaciones que se revisaron para el período 2015-2022, categorizándolas según la etapa del ciclo de desarrollo del ABM en la que se aplicó el ML. Como era de esperarse, las 24 publicaciones sobre aprendizaje por refuerzo (20 %) se utilizan para el aprendizaje en línea en la etapa de exploración, mientras que el aprendizaje no supervisado ha sido usado para el aprendizaje fuera de línea, principalmente en las etapas de validación y análisis.

El aprendizaje supervisado es el más versátil, siendo aplicado incluso en casos en línea durante la etapa de exploración. El aprendizaje fuera de línea incluye 96 publicaciones (80 %). De ellas, hay 36 casos de aprendizaje a priori (37.5 %) que asisten al modelador en la etapa

Tabla 3.3: Artículos sobre ML en ABM revisados por etapa de modelado.

Etapa ABM	Referencias
Diseño	Rosés et al. (2021), Ghaffarian et al. (2021), Wang et al. (2021), García-Magariño et al. (2020), Yao et al. (2020), Pouladi et al. (2020), Augustijn et al. (2020), Abdulkareem et al. (2019), Lee y Hong (2019), Batata et al. (2019), Candadai et al. (2019), Jäger (2019), Hyun et al. (2019), Saeedi (2018), Abdulkareem et al. (2018), Kavak et al. (2018), Sánchez-Maróño et al. (2017), Romero-Mujalli et al. (2017), Sadeghipour et al. (2017), Zhang et al. (2016), Lee y Wong (2016), Sánchez-Maróño et al. (2015), Wolf et al. (2015), Balasubramaniam y Thigarasu (2015)
Exploración	Jäger (2021), Ibrahim et al. (2021), Xu et al. (2019), Harati et al. (2021), Cummings y Crooks (2020), Schauder et al. (2020), Sert et al. (2020), Fuller et al. (2020), Liang et al. (2020), Aghaie y Heidary (2019), Norman et al. (2018), Dehghanpour et al. (2018), Esmaeili Aliabadi et al. (2017), Jalalimanesh et al. (2017b), Jalalimanesh et al. (2017a), Dehghanpour et al. (2016), Salle (2015), Lakić et al. (2015)
Validación	Kim et al. (2021), Ye et al. (2021), Cockrell y An (2021), Zhang et al. (2020), Neri (2019), van der Hoog (2019), Pearce y Slade (2018), Lamperti et al. (2018), Chu et al. (2018), Fievet y Sornette (2018)
Análisis	Shiono (2021), Nunes et al. (2021), Gursoy y Badur (2021), ten Broeke et al. (2021), Xiao y Liu (2021), Chen et al. (2021), Larie et al. (2021), Xu et al. (2020), Koda et al. (2020), R. Vahdati et al. (2019), Edali y Yücel (2019), Janssen et al. (2019), Garg et al. (2019), Perry y O'Sullivan (2018), Ozik et al. (2018), Zhang et al. (2018), Hayashi et al. (2016), Papadopoulos y Azar (2016), Rajabi et al. (2016)

Fuente: Elaboración propia.

de diseño. En la mayoría de estos casos, el modelo aprendido es utilizado por los agentes al tomar decisiones. En otros casos, el modelo aprendido se utiliza para la generación de datos sintéticos. El aprendizaje a posteriori incluye 60 publicaciones (62.5 %), principalmente para calibración, por ejemplo, buscando un subconjunto de parámetros de entrada que pueda reproducir la salida deseada, y para propósitos de análisis, por ejemplo, encontrando la relación entre parámetros de entrada y salida.

A medida que hemos explorado de manera general los usos y aplicaciones del ML en el contexto del ABM, es fundamental desglosar más detalladamente cómo estas técnicas de aprendizaje interactúan con las etapas específicas del proceso de modelado. Específicamente, nos interesa distinguir entre las aplicaciones de aprendizaje que se realizan antes o después de ejecutar una simulación (conocidas como “offline”), y aquellas que operan en tiempo real mientras se desarrolla una simulación (conocidas como “online”). Esta diferenciación es crucial para comprender las capacidades y limitaciones de cada enfoque en el contexto del ABM. En los apartados subsecuentes, se procederá a un análisis riguroso y detallado de cada uno de estos enfoques: el aprendizaje “offline” a priori, el aprendizaje “offline” a posteriori y, finalmente, el aprendizaje “online”

3.2.2.1. Aprendizaje *offline* a priori

El aprendizaje offline a priori se utiliza exclusivamente con fines de diseño, sin embargo, se distinguen 3 tipos de aplicaciones aunque siempre adoptando algoritmos de ML supervisados, incluyendo redes neuronales (10 publicaciones), árboles de decisión (10), regresión logística (5) y redes bayesianas (4). La Tabla 3.4 enumera estos trabajos. Es importante señalar que, dada la naturaleza de este tipo de algoritmo de ML, se asume la existencia de datos del dominio de aplicación.

En un primer grupo, las aplicaciones incluyen la generación de datos sintéticos que serán utilizados posteriormente por el ABM, por ejemplo, [Saeedi \(2018\)](#) genera predicciones de cambio de uso del suelo con redes neuronales e implementa esos resultados como parte del entorno. Otras publicaciones con un propósito similar son [Ghaffarian et al. \(2021\)](#); [Wang et al. \(2021\)](#); [Yao et al. \(2020\)](#).

Por otro lado, existe un conjunto de aplicaciones enfocado en la formulación de las reglas de comportamiento de los agentes. Estas reglas son diseñadas para que el modelador las

Tabla 3.4: Publicaciones que usan ML en ABM para asistir el proceso de diseño por el algoritmo usado.

Algoritmo	Referencias
DT	Rosés et al. (2021), Augustijn et al. (2020), Sánchez-Marño et al. (2017), Sánchez-Marño et al. (2015), Balasubramaniam y Thigarasu (2015)
RF	Wang et al. (2021), Kavak et al. (2018)
XGB	Ghaffarian et al. (2021)
ANN	Batata et al. (2019), Candadai et al. (2019), Jäger (2019), Saeedi (2018), Romero-Mujalli et al. (2017), Wolf et al. (2015)
BN	Abdulkareem et al. (2019), Hyun et al. (2019), Abdulkareem et al. (2018)
LogR	García-Magariño et al. (2020), Lee y Hong (2019), Sadeghipour et al. (2017), Zhang et al. (2016), Lee y Wong (2016)

Fuente: Elaboración propia.

integre en el modelo, como se observa en Augustijn et al. (2020); Sánchez-Marño et al. (2017); Pouladi et al. (2020). Cabe resaltar que para este fin, se prefieren algoritmos que faciliten la interpretación y simplifiquen la derivación de reglas, tales como los árboles de decisión.

Finalmente, el tercer conjunto de aplicaciones busca desarrollar un modelo de ML a partir de datos existentes y posteriormente incorporar dicho modelo en el ABM, poniéndolo a disposición de los agentes. Ejemplos notables son Jäger (2019); Romero-Mujalli et al. (2017); Wolf et al. (2015). A diferencia de los enfoques anteriores, aquí la precisión se prioriza sobre la interpretabilidad. Los algoritmos utilizados en este contexto varían desde redes neuronales hasta bosques aleatorios y otros métodos de ensamble.

3.2.2.2. Aprendizaje *offline* a posteriori

El aprendizaje “offline” a posteriori se utiliza para validación y análisis. Principalmente se emplea el aprendizaje supervisado de ML, pero, a diferencia del caso de aprendizaje a priori, los datos de entrenamiento son generados por el propio ABM. Hay algunas excepciones que utilizan aprendizaje no supervisado, por ejemplo, Peters et al. (2016) combinan un Modelo de Interacción Basado en Plantas con mapas de características autoorganizados (SOM, por sus siglas en inglés), un tipo de red neuronal, para visualizar la relación entre entradas y salidas.

Otro ejemplo es ilustrado por Koda et al. (2020), quienes utilizan el algoritmo de Louvain para analizar la formación de redes sociales y reconstruir estructuras de agrupamiento. Por su

Tabla 3.5: Publicaciones que usan ML en ABM para asistir el proceso de análisis por el algoritmo usado.

Algoritmo	Referencias
DT	R. Vahdati et al. (2019), Hayashi et al. (2016)
RF	Gursoy y Badur (2021), Chen et al. (2021), Edali y Yücel (2019), Garg et al. (2019), Perry y O'Sullivan (2018), Ozik et al. (2018)
SVM	ten Broeke et al. (2021)
ANN	Shiono (2021), Xiao y Liu (2021), Chen et al. (2021), Larie et al. (2021), Xu et al. (2020), Ozik et al. (2018), Zhang et al. (2018)
BN	Rajabi et al. (2016)
LinR	Papadopoulos y Azar (2016)

Fuente: Elaboración propia.

parte [Ye et al. \(2021\)](#), con el propósito de calibrar parámetros para reproducir comportamiento real en sistemas distribuidos, utilizan PCA para comprimir la alta dimensionalidad del espacio de estado del agente. En lo que respecta al trabajo de [Marzouk y Hassan \(2022\)](#), ellos emplean el algoritmo k-means para agrupar el rendimiento de evacuación de visitantes y trabajadores de museos en respuesta a un incendio.

Para la validación, la meta es a menudo encontrar el subconjunto de parámetros que mejor reproduzca un comportamiento deseado. La interpretabilidad del algoritmo, en este contexto, suele ser secundaria. Ejemplos destacados de algoritmos de ML empleados con este fin incluyen algoritmos genéticos ([Cockrell y An, 2021](#); [White et al., 2022](#)), Boost DT ([Zhang et al., 2020](#); [Lamperti et al., 2018](#)), bosques aleatorios ([Pawar y Jha, 2022](#)), máquinas de vectores de soporte ([Perumal y Zyl, 2022](#)), recocido simulado ([Neri, 2019](#)), redes neuronales ([van der Hoog, 2019](#); [Reiker et al., 2021](#); [Jørgensen et al., 2022](#); [Krivorotko et al., 2022](#)), y redes bayesianas ([Chu et al., 2018](#); [Masuda et al., 2022](#)).

En la fase de análisis, los investigadores buscan comprender las relaciones entre las variables de entrada y salida. Sin embargo, un dato interesante es que no se ha dado preferencia a la interpretabilidad o transparencia del modelo, como se muestra en la Tabla 3.5, hay una clara preferencia hacia los algoritmos de redes neuronales (12 artículos) y bosques aleatorios (9). Las aplicaciones incluyen la comprensión de la relación entre parámetros de entrada y salida; por ejemplo, [Janssen et al. \(2019\)](#) proponen analizar la emergencia a través del descubrimiento causal y [Edali y Yücel \(2019\)](#) utilizan bosques aleatorios para entender la dinámica de los sistemas modelados.

También son de interés los modelos sustitutos, donde el ML se utiliza para aproximar el ABM. Por ejemplo, [ten Broeke et al. \(2021\)](#) utilizan máquinas de vectores de soporte para aproximar la simulación del ABM en tres casos, permitiendo distinguir qué parámetros determinan los límites en el espacio de parámetros entre regiones; [Larie et al. \(2021\)](#) emplean una ANN como sistema sustituto para predecir el tiempo de evolución de un ABM y [Papadopoulos y Azar \(2016\)](#) proponen modelos sustitutos con regresiones lineales múltiples.

En cuanto al uso de ML para la validación en ABM, varios estudios han utilizado técnicas de ML para apoyar la validación y calibración de ABM en el campo de la economía. Por ejemplo, [Zhang et al. \(2020\)](#) y [Lamperti et al. \(2018\)](#) utilizaron sustitutos de ML para reproducir el modelo de fijación de precios de activos, mientras que [Neri \(2019\)](#), [Chu et al. \(2018\)](#), y [Fievet y Sornette \(2018\)](#) utilizaron técnicas de ML para aproximar series temporales financieras o económicas en sus ABM.

Para la etapa de análisis, la mayoría de las aplicaciones están concentradas en medicina (7 artículos), ciencias biológicas (6) y ciencias sociales (5). Nuevamente, en el dominio de la medicina, los modelos buscan entender la propagación de enfermedades como el COVID-19 ([Xiao y Liu, 2021](#)), la Leishmaniasis Cutánea ([Rajabi et al., 2016](#)), o la sepsis ([Larie et al., 2021](#)).

En lo que respecta a las ciencias biológicas, las aplicaciones intentan analizar algunos aspectos de las sociedades animales, por ejemplo, formación de grupos ([Koda et al., 2020](#)), mortalidad y paisaje ([Day et al., 2020](#)), comportamiento de búsqueda de alimentos ([Bhattacharjee et al., 2019](#)), regla de autoeliminación en poblaciones de plantas [Ma et al. \(2019\)](#), respuesta a la disponibilidad fluctuante de recursos alimenticios [Scott et al. \(2018\)](#), aptitud reproductiva y personalidad animal [MacPherson et al. \(2017\)](#), y competencia subterránea entre plantas [Peters et al. \(2016\)](#).

Para las ciencias sociales, los problemas abordados están relacionados con el comportamiento o dinámicas humanas, específicamente, fuga de cerebros ([Gursoy y Badur, 2021](#)), dispersión humana en el Pleistoceno tardío en África ([R. Vahdati et al., 2019](#)), y la reconstrucción de interacciones humanas en entornos humanos pasados ([Perry y O'Sullivan, 2018](#)).

3.2.2.3. Aprendizaje *online*

El aprendizaje online se relaciona exclusivamente con la etapa de experimentación, permitiendo a los agentes explorar el entorno para generar una regla de comportamiento óptimo. En este caso, los datos se generan durante la simulación.

En la mayoría de los casos, la comunidad ha adoptado algún algoritmo de aprendizaje por refuerzo, aunque hay instancias en las que se han elegido algoritmos supervisados. Por ejemplo, [Jäger \(2021\)](#) introduce un marco para ABM que utiliza redes neuronales artificiales para aprender reglas de comportamiento, y [Dehghanpour et al. \(2016\)](#) proponen que cada agente, representando a una compañía en el dominio del mercado de energía, desarrolle una red bayesiana dinámica privada del mercado usando aprendizaje bayesiano disperso para entrenar el modelo que luego se usa para inferir el estado futuro del mercado para estimar la función de oferta óptima.

Argumentando que los modelos económicos son demasiado simplistas, [Salle \(2015\)](#) introduce un modelo de expectativas de inflación basado en una red neuronal artificial. Los pesos se inicializan aleatoriamente de manera que, al menos al principio, los modelos mentales de los agentes y las expectativas resultantes difieren. A medida que se disponen de nuevas observaciones, las redes neuronales se entrenan.

El Q-learning es el algoritmo de aprendizaje por refuerzo más utilizado, seguido de otros algoritmos de diferencia temporal. Otras propuestas involucran el uso de Deep Q-Nets, aunque el dominio de aplicación aún está en entornos poco realistas, como el modelo de segregación de Schelling. La mayoría de los trabajos en esta categoría adoptan la regla de comportamiento óptimo aprendida como una hipótesis del fenómeno estudiado.

Sin embargo, tres de ellos asisten al modelador en el diseño del ABM. El enfoque propuesto por [Cummings y Crooks \(2020\)](#) produce agentes que crean su propia estrategia óptima, que el diseñador puede interpretar. [Fuller et al. \(2020\)](#) proponen la aplicación del aprendizaje por refuerzo para permitir a los agentes definir sus reglas de interacción necesarias, y [Norman et al. \(2018\)](#) examinan un nuevo camino para el desarrollo de políticas de agentes en el que los modeladores no elaboran manualmente las políticas, sino que permiten que surjan mediante la aplicación del aprendizaje por refuerzo dentro de un entorno de motor de juego.

Los dominios principales en los que se aplica el aprendizaje online son la energía y

la medicina, con 9 y 3 publicaciones respectivamente. Dentro del sector energético, los mercados eléctricos se distinguen por la presencia de proveedores y consumidores de energía eléctrica. Las investigaciones en este ámbito se centran en variados aspectos, desde las estrategias de suministro hasta el número de agentes modelados.

Un ejemplo de ello es el estudio de [Xu et al. \(2019\)](#), que presenta un mercado de electricidad residencial sensible a la demanda. En este modelo, tanto las empresas generadoras como los minoristas tienen la capacidad de optimizar sus estrategias de suministro mediante el aprendizaje por refuerzo. Por otro lado, [Liang et al. \(2020\)](#) proponen el uso de un algoritmo de gradiente de política determinística profunda. Su objetivo es modelar las tácticas de oferta de las empresas generadoras en escenarios con diversa información, abarcando desde situaciones completamente transparentes hasta las más opacas, y esto se realiza incluso en contextos no discretos.

En cuanto al sector médico, las investigaciones se orientan hacia la identificación de terapias más efectivas. Un claro ejemplo es el trabajo de [Jalalimanesh et al. \(2017b,a\)](#), donde se aplica el Q-learning para mejorar componentes esenciales en radioterapia, incluyendo los esquemas de dosis y de fraccionamiento. De manera similar, la investigación de [Schauder et al. \(2020\)](#) explora la formación de hábitos mediante la diferencia temporal, focalizándose en cómo se establecen preferencias basadas en retroalimentaciones, ya sean positivas o negativas, sobre el consumo reciente de alimentos.

3.3. Discusión

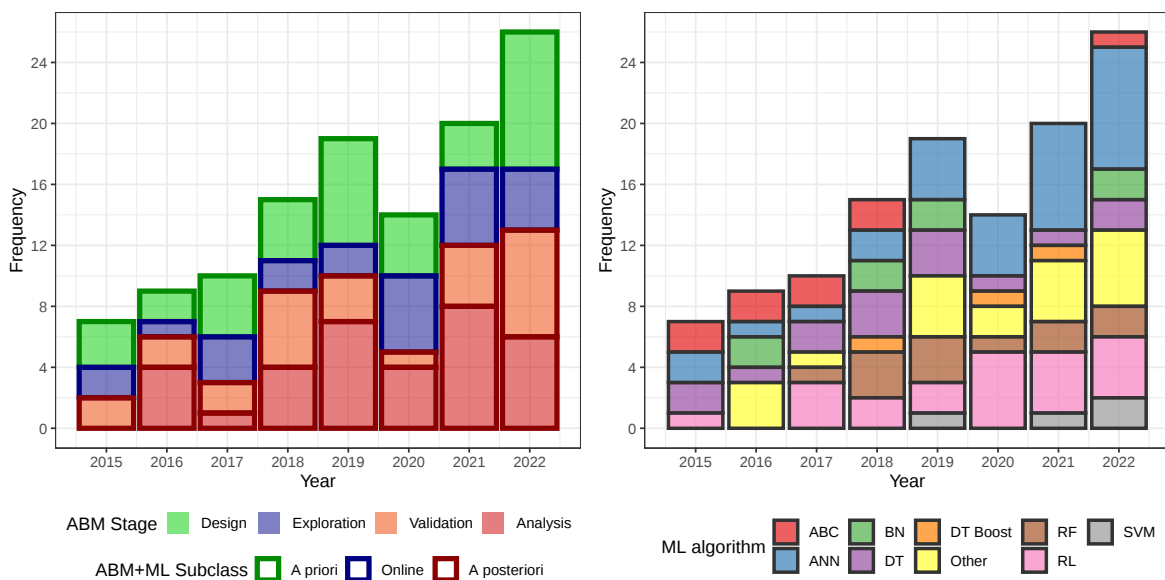
Esta sección discute aspectos clave de la integración de ML en la ABM, abordando tendencias actuales, áreas de oportunidad y desafíos en la documentación. En primer lugar, se analiza la evolución y adopción de diferentes algoritmos de ML en ABM. En segundo lugar, se resaltan enfoques novedosos y poco explorados de ML que podrían impactar significativamente en la modelización de ABM. Finalmente, se aborda la falta de documentación estandarizada para la integración de ML y ABM, proponiendo una extensión del protocolo ODD. En conjunto, se busca orientar futuras investigaciones en este campo en rápido desarrollo.

3.3.1. Análisis de tendencias

Para discutir las tendencias futuras en la aplicación del aprendizaje automático en la modelización basada en agentes, hemos examinado la progresión de las etapas en las que se utiliza el ML y la evolución de las técnicas de ML con el tiempo. La Figura 3.13 resume el número de publicaciones anuales según la etapa del ABM involucrada. La tendencia muestra un número creciente de propuestas, desde siete publicaciones en 2015 hasta 26 en 2022.

Estos resultados confirman la relevancia de estudiar los años recientes para obtener una imagen completa del estado del arte en el campo, lo que concuerda con el análisis exploratorio mostrado en la Figura 3.13. Las etapas de diseño y validación fueron las más respaldadas por ML el último año. El aprendizaje “a posteriori” fuera de línea ha representado entre el 50 % y el 60 % de las publicaciones en los dos últimos años, denotando un interés en las etapas de validación y análisis. En contraste, el interés en el aprendizaje en línea ha disminuido gradualmente del 35.7 % en 2020 al 14.4 % en el último año.

Figura. 3.13: Análisis del número de propuestas revisadas. Propuestas por año y etapas de ABM (izquierda); Propuestas por año y algoritmo de ML (derecha).



Platas-López et al. (2023b).

La Figura 3.13 también proporciona un estudio en profundidad de las tendencias en la adopción de algoritmos de ML. Las ANNs son el algoritmo de aprendizaje supervisado más utilizado, con 29 propuestas. El número de artículos que utilizaron algún algoritmo de ANN

aumentó en más del 100 % durante el período 2020-2021 en comparación con el período 2015-2019. Curiosamente, estos modelos de aprendizaje, considerados cajas negras con poca interpretabilidad, también se están utilizando en la etapa de análisis.

Los algoritmos de aprendizaje basados en metaheurísticas, como los algoritmos genéticos, se están adoptando con más frecuencia, principalmente para tareas de validación o calibración de parámetros. De manera similar, los algoritmos de aprendizaje no supervisado, como k-means y PCA, han comenzado a utilizarse desde 2020, típicamente para analizar resultados de simulación.

En contraste, la adopción de algoritmos con mejor interpretabilidad, como árboles de decisión, redes bayesianas y regresión logística, ha disminuido. Finalmente, el número de propuestas que adoptan aprendizaje por refuerzo también ha aumentado, siendo Q-learning el algoritmo preferido. Sin embargo, actualmente se están explorando algunas propuestas con otros tipos de algoritmos de refuerzo que son capaces de contender con entornos más complejos.

El uso de ML en la metodología de modelización ABM continuará aumentando en los próximos años, debido a la creciente disponibilidad de datos y poder de cómputo, lo que permite el desarrollo de modelos más complejos que pueden ser validados usando técnicas de aprendizaje automático. Además, la creciente adopción de ML en otros campos probablemente inspirará a más investigadores a aplicar estas técnicas a ABM también.

3.3.2. Oportunidades emergentes

Con respecto a nuevas áreas de investigación, varios algoritmos de ML pertinentes aún necesitan exploración, por ejemplo, los algoritmos de aprendizaje no supervisado han sido poco explorados en cualquiera de las etapas de ABM, excepto en la validación, ya sea para agrupamiento o reducción de dimensiones.

El aprendizaje por refuerzo domina las exploraciones para el aprendizaje en línea, a pesar de la pertinencia de los algoritmos supervisados y no supervisados incrementales en este sentido. Ejemplos de estos últimos incluyen tareas de agrupamiento en línea, aprendizaje de representación en línea y detección de anomalías en línea. Una revisión completa de los algoritmos de aprendizaje en línea se puede encontrar en el trabajo de [Hoi et al. \(2021\)](#).

Los algoritmos de agrupamiento en línea podrían detectar patrones emergentes en los resultados de simulación, mientras que los algoritmos de aprendizaje de representación en línea podrían extraer características significativas de datos de alta dimensión. Los algoritmos de detección de anomalías en línea podrían monitorear el comportamiento de los agentes en tiempo real y detectar comportamientos anormales que podrían indicar un problema con el modelo.

Los beneficios potenciales de estos algoritmos incluyen una mayor eficiencia y precisión en la validación y análisis de modelos. Otra área inexplorada es el uso de algoritmos de aprendizaje profundo, como redes neuronales convolucionales, en ABM. Estos algoritmos han mostrado un gran éxito en tareas de procesamiento de imágenes y señales y podrían extraer potencialmente características significativas de los datos de simulación basados en agentes.

Hay varios enfoques de ML novedosos que podrían ser útiles en la modelación de ABM. Uno es el aprendizaje activo, que mejora la eficiencia del proceso de aprendizaje al seleccionar instancias específicas del conjunto de datos para ser etiquetadas. Este enfoque podría reducir el costo de etiquetado en la validación de modelos de ABM, donde a menudo se requiere un gran volumen de datos etiquetados.

Otro enfoque que podría ser útil es el aprendizaje de transferencia, que permite el uso de conocimientos adquiridos previamente en un dominio para mejorar el aprendizaje en otro dominio relacionado. Este enfoque podría aplicarse en la modelación de ABM para transferir conocimientos de un modelo a otro o para usar información adquirida previamente en la construcción de modelos similares.

Otros enfoques de ML novedosos que podrían aplicarse en la modelación de ABM incluyen el aprendizaje por imitación y el aprendizaje auto-supervisado. Estos enfoques podrían abordar diferentes desafíos en la construcción de modelos de ABM, como identificar patrones emergentes y adaptarse a entornos cambiantes.

En general, la aplicación de estos enfoques novedosos de ML en la modelación de ABM podría mejorar la eficiencia y precisión de la construcción de modelos, permitiendo la identificación de patrones y tendencias complejas en los sistemas bajo estudio. Sin embargo, es importante tener en cuenta que estos enfoques pueden requerir un mayor conocimiento

y experiencia en el diseño e implementación de modelos, así como en el procesamiento y análisis de datos.

3.3.3. Estándares de documentación

Otro desafío, aún más crítico, está relacionado con la escasa documentación de las implementaciones. En el peor de los casos, la descripción tanto del algoritmo de ML utilizado como del modelo ABM es mínima. En los mejores casos, se intentó describir ambas herramientas por separado.

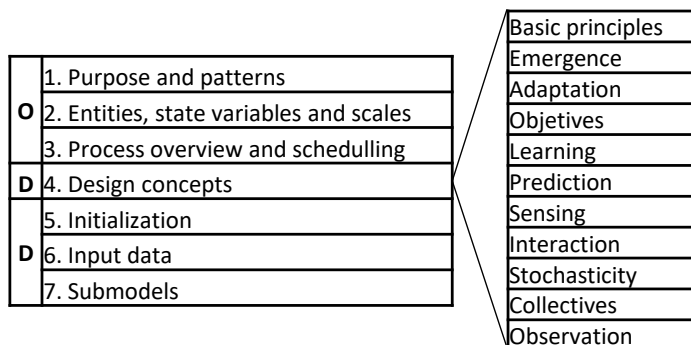
Sin embargo, para ayudar a los desarrolladores menos experimentados en el área, falta una descripción unificada que permita entender la integración de ambas herramientas como una sola. Esto permitiría diseñar un flujo de trabajo claro y transmisible para cada uno de los problemas a abordar.

Proponemos extender el protocolo ODD (Overview, Design concepts, Details) ([Grimm et al., 2010, 2020](#)) para incluir una descripción estandarizada de la integración de ML y ABM. ODD consta de siete elementos como se muestra en la Figura 3.14. Sin embargo, el actual Protocolo ODD puede no ser capaz de acomodar la diversidad de aplicaciones de ML en los ABM.

Afortunadamente, el Protocolo ODD cubre otros elementos necesarios para realizar dichas descripciones, lo que facilita su extensión para documentar las diversas aplicaciones de ML identificadas en la literatura de ABM. Proponemos categorizar las aplicaciones de ML en ABM para habilitar las preguntas de diseño que guían la extensión al componente de aprendizaje del protocolo. Por ejemplo, sugerimos incluir las siguientes preguntas: ¿Cuál es el propósito del aprendizaje? ¿Cuándo se realiza el aprendizaje? ¿Qué componentes del ABM son afectados por el aprendizaje? ¿Cómo se calcula el aprendizaje?

Además, otros métodos de documentación de ABM, como STRESS ([Monks et al., 2018](#)) y RAT-RS ([Achter et al., 2022](#)), también pueden proporcionar valiosos conocimientos para la extensión del Protocolo ODD. Por lo tanto, como trabajo futuro, se propone explorar las posibles contribuciones de estos métodos para asegurarse de que el protocolo extendido sea exhaustivo y cubra todos los aspectos relevantes de la integración de ML y ABM.

Figura. 3.14: Estructura de las descripciones de modelos siguiendo el protocolo ODD.



Fuente: Adaptado de Grimm et al. (2010)

3.4. Conclusión

El uso de métodos de ML en ABM se volverá cada vez más importante, ya que los científicos tienen la intención de adoptar modelos más realistas con mayor complejidad y requisitos computacionales. Inherentemente, las dificultades para diseñar, calibrar y analizar dichos modelos también aumentan, sugiriendo la adopción de métodos de ML para asistir estas tareas.

Este trabajo presenta una revisión novedosa de las últimas tendencias en ABM asistido por ML. La revisión incluye 120 artículos publicados en revistas de alto impacto de 2015 a 2022. Los trabajos se resumieron en función de los problemas que podrían resolverse mediante la aplicación de ML en ABM. Esta revisión permitió mapear los problemas o críticas con las etapas del ciclo de desarrollo de ABM, permitiendo a los modeladores con menos experiencia en este enfoque deducir un flujo de trabajo adecuado al realizar sus implementaciones. Para esto, proponemos la siguiente categorización:

Aprendizaje offline. En sus dos formas:

1. A priori. Antes de que se realice la simulación, se aplica el algoritmo de aprendizaje a los datos existentes del fenómeno que se va a modelar, resolviendo en mayor o menor medida aspectos relacionados con el diseño del modelo. Hay principalmente tres tipos de aplicaciones: instanciar datos en el modelo, derivar reglas de comportamiento para ser implementadas en el modelo y encapsular la simulación dentro del conocimiento disponible de los agentes.

2. A posteriori. Después de que se completa la simulación, se aplican algoritmos de aprendizaje automático a las salidas del modelo. Este enfoque puede abordar dos de las críticas de los ABM: problemas de validación/calibración y problemas de análisis. El primero se refiere a encontrar un subconjunto de parámetros de entrada que generen la salida deseada, mientras que el segundo busca explicar la relación entre entradas y salidas.

Aprendizaje en línea. Durante la simulación, el aprendizaje se aplica a los datos generados por los agentes a medida que interactúan con estructuras de recompensa previamente definidas y estados ambientales para encontrar reglas de comportamiento óptimas.

El uso de métodos de ML en ABM ofrece varias ventajas. Primero, puede ayudar a reducir el costo computacional asociado con la simulación de modelos complejos, permitiendo a los modeladores obtener resultados más realistas y precisos en menos tiempo. Segundo, los métodos de ML pueden ayudar en el diseño y calibración del modelo, que a menudo son tareas difíciles y que consumen mucho tiempo. Tercero, el ML puede ayudar a analizar y entender mejor las relaciones entre las entradas y salidas, así como identificar patrones y comportamientos que pueden no ser inmediatamente obvios.

A pesar de los potenciales beneficios de usar ABM respaldado por ML, persiste un problema significativo en la falta de comprensión de los métodos de ML. Esta falta de conocimiento puede llevar a problemas y omisiones en la presentación de resultados, obstaculizando finalmente el avance de los ABM como herramientas de investigación.

Si bien la implementación de algoritmos de ML es un proceso complejo que involucra varios pasos, algunos investigadores no proporcionan suficientes detalles sobre los pasos tomados para elegir un modelo en particular. Esta falta de transparencia puede llevar a la ambigüedad, lo que socava la credibilidad de la investigación.

Además, el uso de ciertas técnicas de ML, como las redes neuronales, puede estar más impulsado por su popularidad que por su utilidad intrínseca para resolver un problema en particular. Esta tendencia resalta la necesidad de que los investigadores sean más críticos y selectivos en su uso de las técnicas de ML, asegurándose de elegir la técnica más apropiada para su pregunta de investigación específica, en lugar de simplemente seguir la



última tendencia. Al hacerlo, los investigadores pueden mejorar la calidad y precisión de sus resultados y contribuir al avance del campo.

Una crítica significativa a los enfoques existentes para integrar ABM y ML es su incapacidad para proponer métodos generales que se apliquen a muchos problemas y dominios. La falta de un marco metodológico claro y sistemático para implementar estos sistemas compromete su adaptabilidad para enfrentar nuevos problemas. Esta limitación resalta la necesidad de un flujo de trabajo que permita modelos más transparentes y fácilmente replicables, lo que puede proporcionar mayor certeza sobre las ventajas de adoptar el enfoque ABM respaldado por ML.

Desarrollar dicho flujo de trabajo presenta un desafío significativo ya que requiere la integración de múltiples componentes, incluida la adquisición de datos, preprocesamiento, selección de modelos, validación e interpretación. Además, es esencial considerar las características específicas del dominio del problema, incluidos los datos disponibles y el conocimiento existente.

El flujo de trabajo debe ser transparente y fácilmente replicable para permitir que otros investigadores construyan sobre trabajos anteriores y comparen resultados entre estudios. Otro desafío significativo es asegurarse de que el flujo de trabajo sea flexible y adaptable a diferentes dominios de problemas y conjuntos de datos. Esta adaptabilidad requerirá el desarrollo de componentes modulares que puedan intercambiarse fácilmente, dependiendo de las necesidades específicas del problema en cuestión.

3.5. Apéndice



Received: 15 September 2022 | Revised: 14 April 2023 | Accepted: 18 April 2023
DOI: 10.1111/exsy.13325

SURVEY ARTICLE

Expert Systems  WILEY

A survey on agent-based modelling assisted by machine learning

Alejandro Platas-López  | Alejandro Guerra-Hernández |
Marcela Quiroz-Castellanos | Nicandro Cruz-Ramírez

Instituto de Investigaciones en Inteligencia Artificial, Universidad Veracruzana, Veracruz, Mexico

Correspondence

Alejandro Platas-López, Instituto de Investigaciones en Inteligencia Artificial, Universidad Veracruzana, Calle Paseo Lote II, Sección 2a No. 112, Nuevo Xalapa, Xalapa, Veracruz 91097, Mexico.
Email: alejandro.plataslo@anahuac.mx

Funding information

CONACyT, Grant/Award Number: 743662

Abstract

Agent-based models have diversified their applications across various domains due to the ease with which different phenomena can be represented and simulated. These models incorporate heterogeneous, autonomous agents, local interactions, bounded rationality, and often feature explicit spatial representations. However, certain challenges have been identified in their application, including the complexity of design, difficulty in calibrating parameters, and interpreting and analysing results. Therefore, incorporating machine learning (ML) tools in the various stages of the agent-based modelling and simulation process presents a promising approach for current and future research. The main hypothesis of this study is that integrating ML techniques and tools into agent-based modelling can help address challenges encountered during different stages of implementation, ultimately leading to more accurate and effective simulations. The methodology employed in this study involves a comprehensive search and analysis of relevant literature on the topic. This survey reviews significant developments in the integration of ML into the agent-based modelling and simulation process in recent years. The results of this study summarize the fundamental concepts of ML and its applications in agent-based modelling, and provide insights into the prospects and challenges for ML-assisted agent-based modelling in the near future.

KEYWORDS

agent-based modelling, critics, development stages, machine-learning, neural networks, reinforcement learning

1 | INTRODUCTION

Modelling is a fundamental practice in science (Schwarz et al., 2009). A scientific model is an abstract and simplified representation of a system, or phenomenon that makes its core features explicit and visible for generating explanations and predictions. Scientists from different disciplines have resorted to various modelling alternatives. Among them, according to Gilbert (2020), the agent-based modelling (ABM) approach has become popular since it allows individual heterogeneous agents and their interactions to be directly represented by explicitly describing their rules of behaviour in a space. The resulting models enable multiple scales of analysis in a natural way, for example, the emergence of macro or societal-level structures from individual actions, various types of adaptation, and learning. Neither of these features is easily obtained with other modelling approaches, offering much freedom for modelling.

Abbreviations: ABM, agent-based modelling; ML, machine learning.

Expert Systems. 2023;e13325.
<https://doi.org/10.1111/exsy.13325>

wileyonlinelibrary.com/journal/exsy

© 2023 John Wiley & Sons Ltd. | 1 of 23

Indeed, the use of ABM has increased since the beginning of the current century, as shown by the search in Scopus plotted in Figure 1. The number of publications with titles related to this modelling approach (agent-based or individual-based modelling) has increased over the years, resulting in a total of 8875 works, including 4608 journal articles (51.92%), 3462 conference papers (39.01%), 453 book chapters (5.1%), and 179 literature reviews (2.02%). Furthermore, the domain of application of these publications has diversified, as shown in Figure 2: 1485 works have been published in Social Sciences (18.3%), 1251 in Environmental Sciences (15.4%), 653 in Agricultural and Biological Sciences (8%), 626 in Business, Management, and Accounting (7.7%), 511 in Economics, Econometrics, and Finance (6.3%), among others.

Despite this favourable context, some issues of the ABM approach have been identified in the literature, including: its high computational cost; the intricacies involved in its design; the difficulty in calibrating its parameters; the laborious analysis and interpretation of their results; and, the required programming skills for its implementation. These issues affect different stages in the development cycle of simulation methods such as ABM. These stages include design, experimentation, calibration, validation, and analysis, which are shown in Figure 3. It seems that machine learning (ML) algorithms are being used to mitigate such effects, possibly encouraged by increasing availability of data, methods, and computational resources (Russell et al., 2020). The application of ML techniques in ABM has been observed before, during, and/or after running the

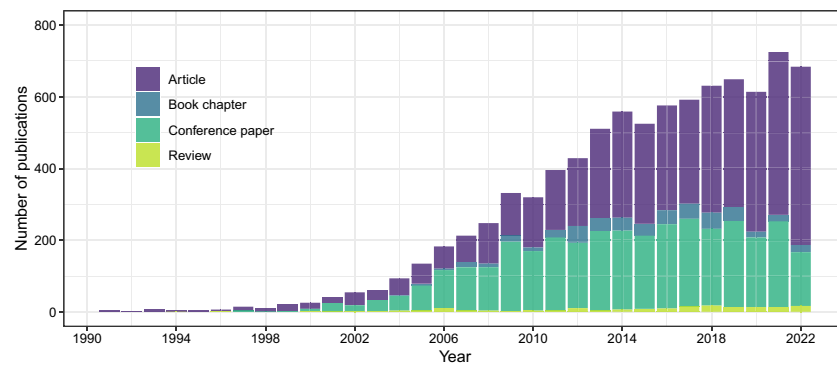


FIGURE 1 Increase in the number of publications related to agent-based and individual-based modelling. Since 2000, the average annual growth rate of published works has been 12.8%. Source: Scopus

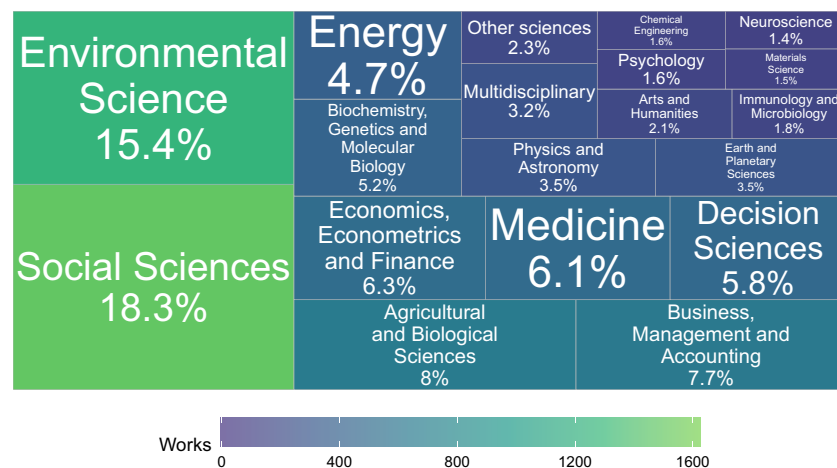


FIGURE 2 Proportion of publications related to agent-based and individual-based modelling by the domain of application. Source: Scopus

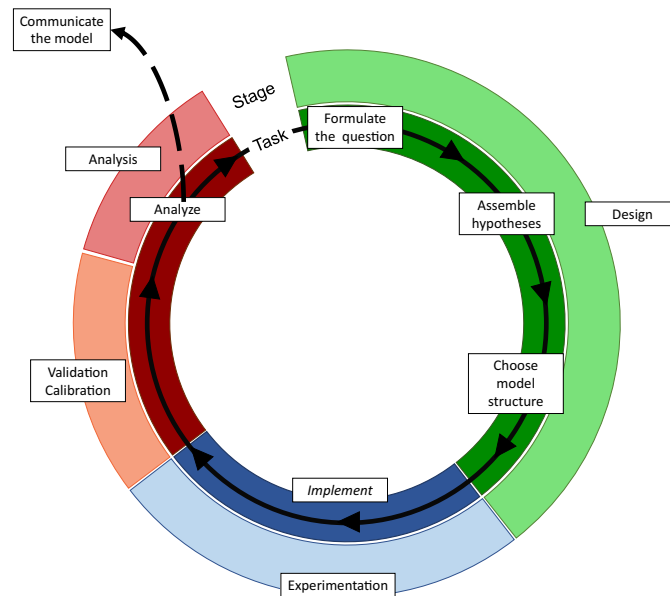


FIGURE 3 The tasks of the agent-based modelling development cycle proposed by Railsback and Grimm (2019), organized in stages

simulation itself, serving different purposes associated with the mentioned development stages, for example, synthesizing agents from existing data, enabling adaptive agents during simulation, calibrating parameters, analysing outputs, and so forth.

Although the number of publications exploring the synergy between ML and ABM has increased in recent years, these efforts are still scarce and somewhat disorganized. Quite often, they constitute isolated contributions addressing particular problems instead of attending to the emerging methodological issues resulting from the integration of these techniques. Therefore, our objective is to provide a broad but integrated picture of the current state of knowledge about the ML and ABM interaction, focusing on the types of ML algorithms being used at different moments in the ABM development cycle to cope with the current limitations of the approach. The emerging methodological issues in these efforts are also of our interest. The main contributions of this literature review are as follows:

1. This study provides an exhaustive review of 120 published works that apply ML techniques to ABM within the period of 2015–2022, covering more than two-thirds of the publications during the observed period. This coverage allows for a comprehensive and representative overview of the advances and trends in this research area. To ensure the quality of the review, strict quality assessment criteria were used in the selection of the works. Furthermore, a taxonomy based on the moment in which ML is applied and the involved ABM development stages is proposed:

Offline learning: ML is not applied during the simulation:

- a. *A priori:* This case is related to the design and calibration stages of the ABM development cycle. ML is applied to existing data to instantiate model entities and/or calibrate parameters.
- b. *A posteriori:* This case applies to the validation and analysis stages of the ABM development cycle. ML algorithms are applied to data produced by the simulation to explain the outputs and/or validate the behaviour of the system.

Online learning: ML is applied during the simulation in the experimentation or exploration stages of the ABM development cycle. Here, the agents generate their own data in order to obtain findings from the simulated phenomenon.

2. Some areas of opportunity for improvement have been identified, particularly the need to adopt methodologies that enable both the ML and ABM components to be easily understood and replicated within concrete and transparent workflows. Although standardized documents such

as the Overview, Design concepts, Details (ODD) protocol (Grimm et al., 2010, 2020), STRESS (Monks et al., 2018), and RAT-RS (Achter et al., 2022) provide a consistent, logical, and readable account of the structure and dynamics of ABM, their coverage of learning is limited. Recently, proposals to extend the ODD protocol to incorporate ML into the ABM workflow have been put forth by Platas-López et al. (2023). However, it would be pertinent to investigate possible extensions to STRESS and RAT-RS too.

The rest of the survey is structured as follows: Section 2 provides context by characterizing the criticisms of ABM reported in the literature, identifying the kind of ML algorithms considered to face them, and comparing the approach followed in this survey with other ones of interest. Section 3 presents the adopted methodology and the articles included in the literature search performed for this work. Section 4 presents the results of the bibliographic analysis for understanding the thematic development of the area. In Section 5, the results of the literature review are presented. Section 6 discusses the trends of the latest works and the most significant challenges when integrating ML and ABM. Finally, Section 7 offers conclusions.

2 | FOUNDATIONS AND PRIOR INVESTIGATIONS

In order to guide this literature review by general principles, we decided to explore the criticisms of ABM reported in the literature to identify opportunity areas where ML has been applied to face them. In what follows, these criticisms are identified and related to the ABM stages they affect. Then, the pertinence of different kinds of ML algorithms is established by recalling their usual categorization in supervised, unsupervised, and reinforcement methods. Finally, the approach followed in this survey is compared with other surveys of interest.

2.1 | Criticisms to the ABM approach

A search for criticisms of ABM was conducted on Google Scholar, Web of Science, and Scopus. Keywords related to ‘criticisms of agent-based models’, ‘problems with agent-based models’, and other relevant terms were used. The period of interest was established from 2000 to the present, coinciding with the increase of interest in ABM established before (see Figure 1). The selected authors have made significant criticisms of ABM, but also have published relevant works in different scientific disciplines, such as economics, social sciences, computer science, and natural sciences, providing a broader multidisciplinary perspective on the criticisms of ABM. The main criticisms of ABM found in the literature can be categorized as shown in Table 1, including:

Computational cost: ABM has a high computational cost, due to the need for executing several runs of the simulations to generate conclusions that are statistically significant. Furthermore, if the number of agents in the model increases, the execution time and memory requirements also increase.

Design issues: Some ad hoc models, with a high degree of complexity, lack the level of detail appropriate to the phenomenon that is intended to be reflected. There is also a concern about whether the models are correct since sometimes the behaviour mechanisms are not properly understood.

Parameter calibration: Some of the parameters could be sensitive to small variations, that is, a small change at the beginning of a run could produce a totally different output. Indeed, there is a lack of clear indicators for describing the uncertainty dynamics of the input–output process. This is problematic for establishing a precise and accurate understanding of how the inputs influence the outputs, particularly if the level of detail chosen for the model is wrong, as this can lead to unnecessarily complex models with many parameters.

Interpretation: The patterns and results of the interactions among the agents are intrinsically unpredictable, and it is not easy to determine which macro-process is the consequence of a certain micro-process at the agent level. These issues impact visualization, statistical issues related to defining the number of runs and testing hypotheses, spatial exploration of solutions, and sensitivity analysis, as well as effective communication to non-technical audiences.

Programming skills: This has to do with the technical requirements to implement a model in a programming language. In turn, it is also related to the level of complexity intended for the model; it is clear that a simple model will be less difficult to programme than one with many processes and types of agents.

Given that ML also has strong programming skill requirements, this work will focus on the rest of the criticisms, that is, computational cost, design issues, parameter calibration, and interpretation. These criticisms interact in different ways through the stages of the ABM development cycle. For instance, the flexibility offered by ABM is usually considered an advantage, but sometimes it can lead to unnecessarily more complex models in the design stage, depending on the developer's skills. If the model grows in complexity, the number of parameters tends to increase, affecting the parameter calibration stage and increasing the computational cost of the exploration and experimentation stage. The validation and analysis stage is also affected by the complexity, obscuring the relationship between the behaviour of the agents and their effects at the system level, leading to poor interpretations of the simulations. As pointed out by Siegfried (2014)—a bad decision in the level of detail in the model can

TABLE 1 Main criticisms of agent-based modelling in the literature

Authors	Critics				
	Computational cost	Design	Calibration	Interpretation	Programming skills
Axtell (2000)	✓				
Charteris et al. (2001)	✓	✓			
Breckling (2002)	✓	✓	✓		
Tesfatsion (2002)					✓
Li et al. (2008)				✓	
Nourqolipour and Shariff (2010)	✓				
Ligmann and Sun (2010)		✓	✓		
Malleson et al. (2010)		✓			✓
Rand and Rust (2011)	✓				
Miller et al. (2012)	✓				
Lengnick (2013)		✓		✓	
Conte and Paolucci (2014)	✓	✓			
Ghorbani et al. (2014)		✓		✓	
Kostadinov et al. (2014)	✓				
Siegfried (2014)	✓	✓	✓	✓	
Baldwin et al. (2015)				✓	✓
Lee et al. (2015)				✓	
Wilensky and Rand (2015)	✓				
Tarvid (2016)	✓	✓		✓	
Broniec et al. (2021)	✓				

trigger models with high computational consumption and possibly with erroneous calibrations. Figure 4 shows these relationships between the main criticisms of ABM found in literature and the stages of their development cycle.

2.2 | Machine learning algorithms

This section provides a brief introduction to ML, aiming to describe the different tasks that can be solved with this kind of algorithms. There are three types of feedback, provided as input data, that determine the three main kinds of ML algorithms (Russell et al., 2020; Witten et al., 2017): supervised, unsupervised, and reinforcement learning.

2.2.1 | Supervised learning

In supervised learning, the ML algorithm receives as feedback a set of training examples in the form of input-output pairs and learns a function that maps from input to output. The input is usually known as the attributes of the example and the output as its label. Both the attributes and the label can be qualitative, also known as categorical, discrete, or factor variables; or quantitative. This distinction leads to a naming convention for prediction tasks: regression when quantitative outputs are predicted and classification for qualitative ones. Formally, supervised learning is defined as follows: Given a training set of N examples $\{(x_1, y_1), (x_2, y_2), \dots, (x_N, y_N)\}$ where each pair was generated by an unknown function $y = f(x)$, discover a function h that best fits the true function f . The function h is known as the hypothesis and it is drawn from a hypothesis space Russell et al. (2020). It is expected that the learned function h generalizes enough to assign an appropriate label to new examples, beyond the training ones.

Indeed, the real performance measure of an ML algorithm is not how well it learned the training examples, but how well the hypothesis h can handle input data that it has not seen yet. Therefore, a second set of examples (x_i, y_i) , called the test set, is used to measure the accuracy of h . The hypothesis is said to generalize well if it correctly predicts the outputs y of a test set. As can be seen, in this learning approach it is imperative to have sufficient and high-quality labelled examples in order to form good hypotheses of the phenomenon under analysis. Some types of

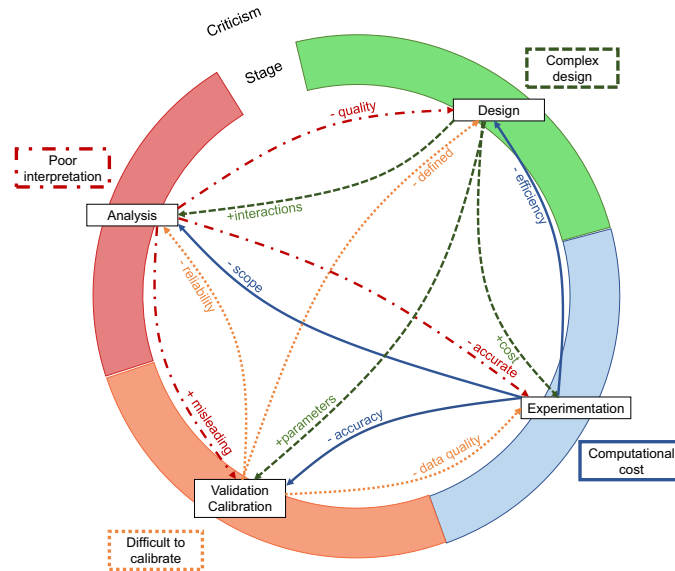


FIGURE 4 Main criticisms towards agent-based modelling (ABM) and the stages in the ABM development cycle they affect

algorithms used for supervised learning are decision trees (DT; Breiman, 2001; Quinlan, 1986), Bayesian networks (BN; Friedman et al., 1997), linear regressions (LR; Montgomery et al., 2015), support vector machines (SVM; Hearst et al., 1998), artificial neural networks (ANN; Haykin et al., 1999), and gradient boosting (GB; Friedman, 2001).

2.2.2 | Unsupervised learning

In unsupervised learning, the ML algorithm searches for patterns in the input without explicit feedback. The most common unsupervised learning task is clustering, that is, detecting similar groups within the input data. Unlike the supervised approach, in this case, we have a set of N observations $(x_1, x_2, \dots, x_N)$ of a p -random vector X with a joint density $Pr(X)$. The goal is to directly infer the properties (patterns) of this probability density without the help of a supervisor or teacher who provides the correct answers or degree of error for each observation.

With supervised learning there is a clear measure of success or lack thereof, that can be used to judge the appropriateness in particular situations and to compare the effectiveness of different methods in various situations. In the context of unsupervised learning, there is no such direct measure of success. It is difficult to determine the validity of the inferences drawn from the results of most unsupervised learning algorithms. Heuristic arguments must be used not only to motivate the algorithms, as is often the case in supervised learning but also to judge the quality of the results. This uncomfortable situation has caused a strong proliferation of methods since the effectiveness is a matter of opinion and cannot be verified directly. Common unsupervised learning algorithms include principal component analysis (Abdi & Williams, 2010) and K-means algorithm (Navarro et al., 1997; Sculley, 2010).

2.2.3 | Reinforcement learning

Unlike the previous learning approaches, reinforcement learning algorithms do not learn passively from a received training set, since in some real applications the examples needed to learn simply do not exist (Russell et al., 2020) and the feedback must be actively collected by the learner. In this approach, the ML algorithm learns from the interaction with the environment (Sutton & Barto, 2018). Reinforcement learning can be viewed as the study of planning and learning in a scenario where a learner actively interacts with the environment to achieve a given goal. This active

interaction justifies the use of the term agent to refer to the learner. The agent's goal achievement is typically measured by the reward it receives from the environment and seeks to maximize.

Reinforcement learning is performed using Markov decision processes (MDPs). An MDP is defined by a set of states S , a set of actions A , a state transition function T , and a reward function R . At each time step, the agent takes an action a_t in state s_t , which produces a reward r_t and a new state s_{t+1} according to the state transition function. The goal is to learn a policy π that maximizes the expected long-term reward. However, the agent only receives immediate rewards from the environment and does not receive feedback on future or long-term rewards. Therefore, it is important to consider delayed rewards or penalties in reinforcement learning. This creates a trade-off between exploring unknown states and actions to gain more information about the environment and rewards and exploiting the information already gathered to optimize the reward. This trade-off is known as the exploration versus exploitation trade-off, which is inherent in reinforcement learning. Some of the most well-known algorithms are Q-learning (Watkins, 1989), and SARSA (Rummery & Niranjan, 1994).

2.3 | Comparison with other surveys

Given the rising interest in the application of ML methods in ABM, some surveys have been published with different particular goals. For instance, Ashreeta et al. (2019) focuses on the energy markets domain. Their analysis of 18 works published between 2001 and 2019 distinguishes three kinds of applications: the use of ML to forecast external input data for the ABM; the use of ML algorithms to implement the learning behaviour of agents when they place bids on the market, and; the use of reinforcement learning for intelligent bidding. Heppenstall et al. (2020) focus on geographical aspects of ABM, arguing that ML contributes to fully broadening and deepening the understanding of geographical systems for developing dynamic spatial simulations. Particularly, they state that the recent interest in ML might offer tools better suited to interrogate the new data that is currently being generated, with applications in behaviour, validation, and emulation.

The review by Zhang et al. (2021) is broader than the previous ones, covering 104 papers published between 1998 and 2021. They identify some applications at the agent (micro) and model (macro) levels. At the micro level, applications include agent situational awareness learning, in which ML is used to predict social, economic, and environmental variables or agent-related behaviours; and a sort of data mining based on the idea that, while acting in their environment, agents can gather data useful for inferring decision-making models by applying appropriate ML techniques. At the macro level, applications include the use of ML to build the mapping structure (emulator) between the parameters and the emergency exit of the ABM, which is particularly useful when the number of agents and parameters grows, generating a high computational cost; and the use of ML to enhance the decision-making of the policymaker, where the policymaker acts as a single macro agent. Finally, they consider some important challenges in the application of ML methods in ABM: obtaining data with sufficient quality to have robust ABM systems in real-world problems; coping with the fact that some models obtained with ML methods do not provide a clear explanation of their predictions, behaving as black boxes; and finding the best method to incorporate bounded rationality in algorithms.

Unlike the work of Ashreeta et al. (2019) and Heppenstall et al. (2020), the present review offers a broader multidisciplinary vision of the synergy between ML and ABM. Our review is guided by the application of ML to cope with some of the disadvantages of ABM reported in the literature, independently of the application domain. Consequently, unlike the work of Zhang et al. (2021), these applications are mapped to the stages of the ABM development cycle, distinguishing if the learning takes place offline or online. This categorization is useful because it allows less experienced researchers to identify the stages of the ABM development cycle that can be assisted by ML. We analyse some of the significant challenges in the area and discuss a framework for incorporating ML in ABM that facilitates understanding, transmitting, and replicating the models.

3 | METHODOLOGY

As previously mentioned, the main goal of this literature review is to present a comprehensive overview of the current state of knowledge on the application of ML methods in ABM. Emphasis is placed on the types of ML methods used at different stages of the ABM development cycle, as well as any adopted integrative approach, if one exists. This objective is achieved through the synthesis of reviewed literature, which highlights various methodological approaches and relevant studies.

The focus of this review is not only on research practices and applications of ML in ABM but also on the adopted methodology. This includes documentation strategies, model design, calibration, validation, analysis, and interpretation of resulting data, as well as methodological strengths and weaknesses. By doing so, this review identifies the primary problems, trends, complexities, and controversies at the heart of the study topic. It also identifies possible directions for future research, the problems that must be explored and addressed, and the application domains of interest.

An important distinction is made based on the moment when ML is applied in the ABM development cycle. Online learning refers to cases where ML is performed during simulation. Although any type of ML can be used for this, it is natural to conceive it in terms of agents exploring their environment to maximize their performance, an instance of reinforcement learning. In contrast, offline learning refers to cases where ML is

performed before (a priori) or after (a posteriori) running the simulation. A priori learning is used during the design stage, while a posteriori learning is used during the calibration, validation, and analysis stages.

Efron and Ravid (2018) present a continuum of literature reviews that ranges from a scientific-quantitative framework, such as systematic reviews, to an interpretive-qualitative framework, such as hermeneutic-phenomenological reviews. The traditional narrative review falls between these two extremes. This literature review departs from the traditional narrative style and adopts some elements of systematic reviews by revealing the inclusion and exclusion criteria, as well as the formula for searching bibliographic material, for the sake of replicability. Figure 5 summarizes the process of selection, review, and analysis of the literature.

The literature search was performed in Scopus and the Web of Science Core Collection. The search strings related to ABM were run on the titles of the publications, while those related to ML were searched in either the title, abstract, or keywords. It is worth noting that specific algorithms were included in the query because some authors refer to them without explicitly mentioning ML.

Scopus search string

TITLE (("agent-based model" OR "individual-based model") OR ("agent-based simulat" OR "individual-based simulat").

AND TITLE-ABS-KEY ("machine learning" OR "neural network" OR {NN} OR {RNN} OR {ANN} OR "decision tree" OR "random forest" OR "logistic regression" OR "Gradient Boosting" OR "Support Vector Machines" OR {SVM} OR "supervised learning" OR "Nearest Neighbor" OR "Naive Bayes" OR "Bayesian network" OR "unsupervised learning" OR "k-mean" OR "k-median" OR "Association Rules" OR {RL} OR {MARL} OR "Temporal Difference" OR "Q-learning" OR {SARSA}).

Web of science search string

TI = (("agent-based model" OR "individual-based model") OR ("agent-based simulat" OR "individual-based simulat").

AND (TS = ("machine learning" OR "neural network" OR {NN} OR {RNN} OR {ANN} OR "decision tree" OR "random forest" OR "logistic regression" OR "Gradient Boosting" OR "Support Vector Machines" OR {SVM} OR "supervised learning" OR "Nearest Neighbor" OR "Naive Bayes" OR "bayesian network" OR "unsupervised learning" OR "k-mean" OR "k-median" OR "Association Rules" OR {RL} OR {MARL} OR "Temporal Difference" OR "Q-learning" OR {SARSA})).

An initial search yielded 340 records in Scopus and 170 in the Web of Science. After eliminating duplicates, 380 articles, conference proceedings, or book chapters were identified. These records were exported in CSV format, which included essential information such as titles, author names and affiliations, abstracts, keywords, and references.

At this point, a bibliometric analysis was conducted with the 380 articles to establish a general overview of the available literature and help identify the main themes. Subsequently, exclusion criteria were applied to ensure that only relevant and high-quality studies were included. By applying these criteria, studies that do not meet the objectives of the literature review or that present significant methodological limitations were not considered. The quality assessment criteria used to guide the search included:

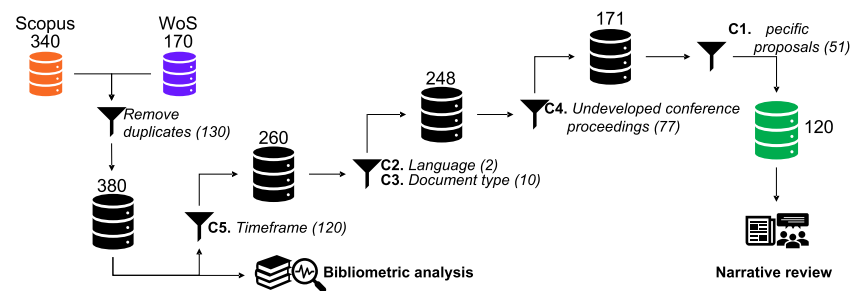


FIGURE 5 Process of selection and review of the literature on ML in ABM. ABM, agent-based modelling; ML, machine learning

C1 To remove all publications that do not present any type of specific proposal related to ABM problems solved from an ML approach with associated experimentation.

C2 To limit international publications written in the English language.

C3 To exclude book chapters and editorials.

C4 To exclude conference proceedings from the initial search, but the complementary search can add specific works if they present a sufficiently developed work.

C5 To restrict publications to the time frame of 2015–2022 to ensure that they are based on the latest and most relevant research within the field.

With criterion C5, 120 articles published before 2015 are removed. Continuing with the selection of works, by applying criteria C2 and C3, 12 articles are excluded. By applying C4, 85 works are initially removed, but eight are reconsidered for their contributions, one of which is a literature review (Ashreeta et al., 2019). Finally, a comprehensive review of the remaining 171 works is performed following the C1 criterion, resulting in 120 works being kept from this database.

4 | BIBLIOMETRIC ANALYSIS

Bibliometric analysis is a quantitative evaluation of published scientific articles, books, or book chapters (Pritchard, 1969). Unlike systematic literature reviews, bibliometric reviews are suitable for domains characterized by large amounts of bibliographic information (Goodell et al., 2021). Therefore, a bibliometric analysis will be used to refine the inclusion/exclusion criteria. The records found in the search were published between 1999 and 2022. Figure 6 shows an annual growth rate of 15.41% in the amount of bibliographic material published on ABM and ML. If the same trend persists, it is estimated with a generalized additive model that between 73 and 78 articles, chapters, or books will be published on this subject by the year 2026.

In bibliometric analysis, a thematic map represents the relationships between key terms in a set of documents, using different colours to display the frequency or co-occurrence of the terms. It consists of an analysis of word co-occurrence networks to define what the research field is discussing, main topics, and trends. Interpreting a thematic map based on density and centrality involves identifying the most relevant or important terms in the set of documents, as well as the relationships between them. Density refers to the frequency with which key terms appear in the documents, while centrality refers to the relative position of terms in the network of relationships between them. The most central terms are those that are more connected with other terms, suggesting greater relevance or importance in the set of documents. There are four types of themes that refer to different categories of topics or research areas in a specific field:

Niche themes: They are characterized by high density and low centrality. This means that there are many articles or publications on a specific topic, but they are not frequently cited or connected to other topics in the field.

Motor themes: They have high density and high centrality. This means that there are many articles or publications on a specific topic, and they are frequently cited or connected to other topics in the field. Motor themes are often considered the driving force behind research in a field.

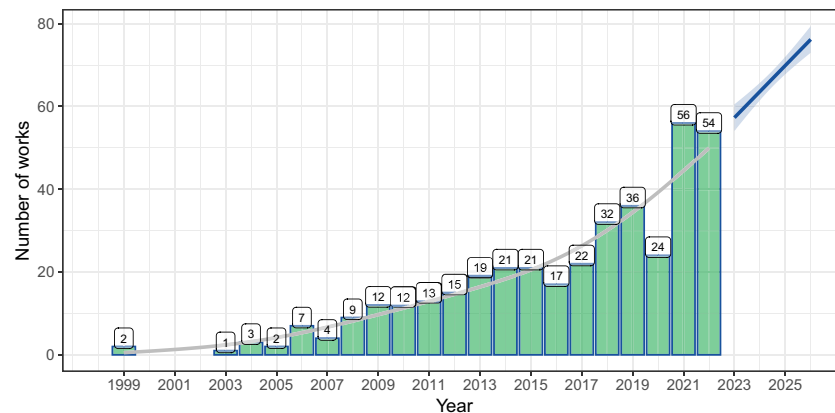


FIGURE 6 Publication trend on agent-based modelling and simulation and machine learning

Emerging or declining themes: They have low density and low centrality. This means that there are relatively few articles or publications on a specific topic, and those that exist are infrequently cited or connected to other topics in the field. Emerging themes are those that are gaining importance and attention in the field while declining themes are those that were once important but are now losing relevance.

Base themes: They have low density and high centrality. This means that there are relatively few articles or publications on a specific topic, but they are frequently cited or connected to other topics in the field. Base themes are often considered foundational concepts or theories in a field.

The thematic map for the selected literature on ML applied to ABM is displayed in Figure 7. It was computed using the key words provided by the authors because they capture the content of the articles with greater variety and precision than the standardized keywords provided by the publishers, which focus on fewer topics related to computation and simulation, without differences in the area's structure. It is worth noting that the author's keywords were manually lemmatized to reduce them to their base form, standardize them, and make them easier to analyse and compare. For instance, common expressions or algorithms were reduced to their abbreviated form, such as ABM, ML, or RL for reinforcement learning.

Niche themes in the upper left quadrant are also known as highly developed and isolated themes. They have well-developed internal linkages (high density) but almost no external linkages and therefore are of limited importance to the field (low centrality). Reinforcement learning, Q-learning, electricity market, game theory, emergence, and supply chain management are isolated issues that have been extensively studied in specific fields. However, integrating them into a comprehensive framework is still challenging.

Emerging or declining themes in the lower left quadrant have low centrality and density, meaning that they are weakly developed and are marginal or with little production. We can see that the references to multi-agent systems (MAS) and land-use change are on the decline. However, it is important to note that research on ABM and MAS has been progressing independently, even if they are closely related. On the other hand, emerging concepts such as data mining, Bayesian networks, and random forest algorithms are gaining increasing attention in the literature.

Basic themes in the lower right quadrant are also known as core or cross-cutting themes. They are characterized by high centrality and low density. These general themes are important to a field of research and cut across the different areas of the field. Based on the query search terms, it is evident that the topics that stand out in the basic themes quadrant are related to ABM as well as ML, indicating that they are highly connected to other topics and appear frequently in the literature. However, it is worth noting that other terms such as artificial neural networks, genetic algorithms, decision-making, and optimization also appear in the quadrant, albeit with lower frequency. These topics may not have as many direct connections with other topics in the network, resulting in lower density. Nonetheless, they are still relevant to the overall theme and suggest that there is ongoing research and exploration in related areas, which could lead to further developments and advancements in these fields.

Motor themes in the upper right quadrant are characterized by their high centrality and density. This means that they are developed and important to the field of research. Among the most developed motor topics in the literature, the approximate Bayesian computation (ABC) method, used for parameter estimation, has been of particular interest in the biological area, where the term individual-based modelling (IBM; Grimm, 1999; Judson, 1994) was coined. In general, it is possible to notice that the interest is related to offline learning applications applied after simulation, in calibration and analysis tasks.

To understand whether the topics of ABM and ML interact with other terms and mainly the way in which such terms interact, a joint word analysis is used through the network of the keywords provided by the authors. Figure 8 shows the conceptual structure among the most

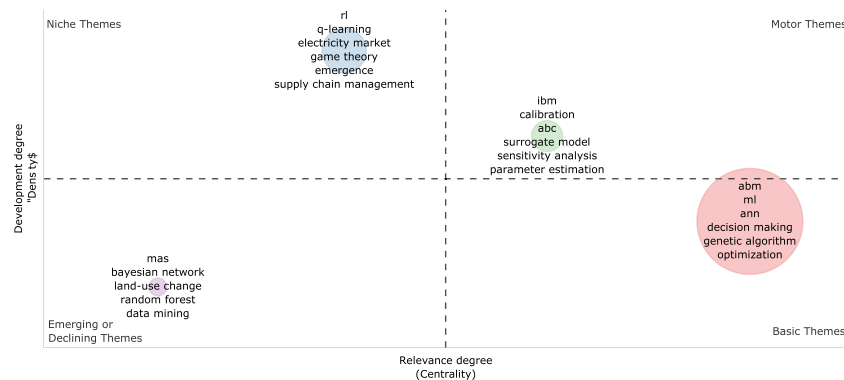


FIGURE 7 The thematic map for the selected literature on ML applied to ABM. ABM, agent-based modelling; ML, machine learning

recurrent problems addressed by researchers in the area. From this analysis, the ML algorithms used in the area, such as artificial neural networks, approximate Bayesian computation, genetic algorithms, Bayesian networks, random forest, and Q-learning.

Each of these ML algorithms is linked to specific stages of the ABM development cycle and to specific application domains. For instance, the approximate Bayesian computation algorithm has a close relationship with the calibration stage and sensitivity analysis, which suggests a need to understand the effect of parameters on their respective outputs. It is also observed that this technique is more frequently used in biological sciences due to its relationship with the term individual-based modelling. Another example related to reinforcement learning algorithms, specifically Q-learning, is their use for optimization purposes for trading in different markets.

A thematic evolution map is a visualization tool that shows the development of a research area over time in terms of changes in its subjects of interest. It is also computed using the lemmatized keywords provided by the authors. Identified themes in the co-occurrence of keywords are plotted with the x-axis representing time and the y-axis representing the relevance or frequency of the themes. These maps provide a useful visualization tool for researchers and practitioners looking to gain a better understanding of the development and trends of a particular research field. By analysing the map, researchers can identify important themes and areas for future investigation, while practitioners can use the map to stay up-to-date on the latest research and developments in their field.

The thematic map for the selected publications on applications of ML in ABM is divided into four periods, as shown in Figure 9. The number of published works and their relative frequency are also shown there. It should be noted that the ABM topic remains throughout the time frame. The rest of the terms may be divided into basic concepts derived from the search query, for example, the involved ML algorithms, and concepts related to the stages of the ABM development cycle and its application domains.

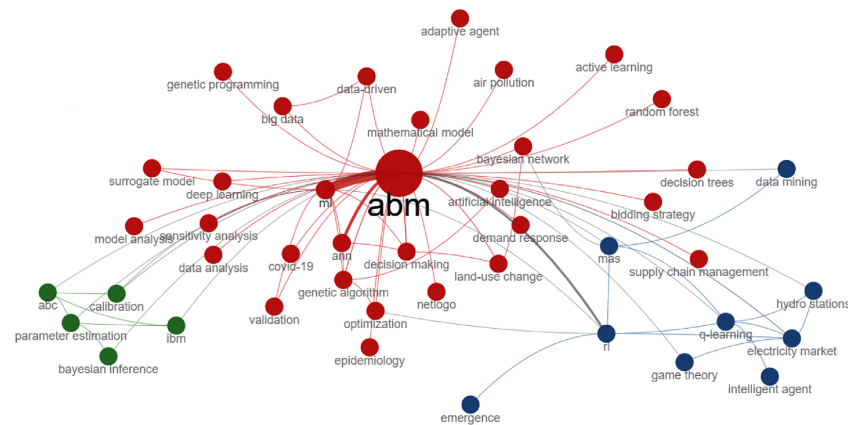


FIGURE 8 Most recurring keywords related to biological sciences (green), reinforcement learning (blue), and others (red). The width of the lines joining the terms is proportional to their frequency of interaction

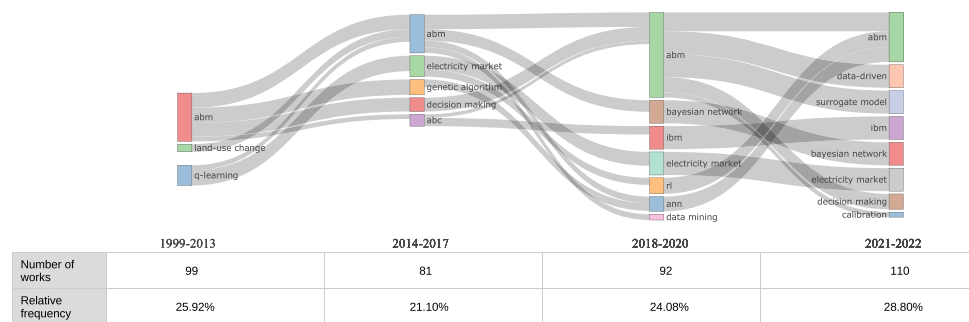


FIGURE 9 A thematic evolution map of topics in the application of ML in ABM. ABM, agent-based modelling; ML, machine learning



In the first period (1999–2013), studies of land-use change were popular, and Q-learning, a reinforcement learning method proposed in 1989 (Watkins, 1989), was quite common. These terms do not appear again in the following periods. In the second period (2014–2017), studies on electric markets and decision-making became popular, while the ML methods commonly adopted were genetic algorithms and approximate Bayesian computation (ABC). The application domain of the electric market is recurrent in the following periods, while decision-making becomes relevant again in the third period (2021–2022). The electric market domain is strongly related to reinforcement learning, which is manifested in the transition to the third period. Approximate Bayesian computation is widely used in individual-based models, as observed in the third and fourth periods. Indeed, there are no emerging application domains in these periods, and the interest seems to be focused on ML algorithms, for example, Bayesian networks and artificial neural networks. In terms of application purposes, the focus has been on data-driven methods for designing models from data and surrogate models for analysis and calibration.

5 | THE APPLICATION OF MACHINE LEARNING IN AGENT-BASED MODELLING

Table 2 lists the 120 publications that were reviewed for the period 2015–2022, categorizing them according to the stage of the ABM development cycle in which ML was applied. Figure 10 complements this categorization by showing the kind of ML algorithms adopted in each case. As expected, the 24 publications on reinforcement learning (20%) are used for online learning in the exploration stage, while unsupervised learning has been used for offline learning, mainly in the validation and analysis stages. Supervised learning is the most versatile, being applied even in online cases during the exploration stage. Offline learning includes 96 publications (80%). Among them, there are 36 cases of a priori learning (37.5%) that assist the modeller in the design stage. In most of these cases, the learned model is used by the agents when making decisions. In other cases, the learned model is used for synthetic data generation. A posteriori learning includes 60 publications (62.5%), mainly for validation, for example, searching for a subset of input parameters that can reproduce the desired output, and for analysis purposes, for example, finding the relationship between input and output parameters.

5.1 | Offline a priori learning

Offline a priori learning is used exclusively for design purposes, always adopting supervised ML algorithms, including neural networks (10 publications), decision trees (10), logistic regression (5), and Bayesian networks (4). Table 3 lists these works. Note that, given the nature of this kind of

TABLE 2 An overview of the reviewed papers (2015–2022), categorized by the stage of the ABM development cycle where ML was applied

ABM stage	References
Design	Taghikhah et al. (2022), Aslani et al. (2022), Catullo et al. (2022), Kadinski et al. (2022), Marvuglia et al. (2022), Salas-Rodríguez and Rivas-Tovar (2022), Zornić and Marković (2022), Liang et al. (2022), Peralta et al. (2022), Rosés et al. (2021), Ghaffarian et al. (2021), Wang et al. (2021), García-Magariño et al. (2020), Yao et al. (2020), Pouladi et al. (2020), Augustijn et al. (2020), Snyder et al. (2019), Abdulkareem et al. (2019), Lee and Hong (2019), Batata et al. (2019), Candadai et al. (2019), Jäger (2019), Hyun et al. (2019), Saeedi (2018), Abdulkareem et al. (2018), Kavak et al. (2018), Bhattacharjee et al. (2018), Cenek and Franklin (2017), Alonso-Betanzos et al. (2017), Romero-Mujalli et al. (2017), Sadeghipour et al. (2017), Zhang et al. (2016), Lee and Wong (2016), Sánchez-Marroño et al. (2015), Wolf et al. (2015), Balasubramaniam and Thigarasu (2015)
Exploration	Park and Ryu (2022), Ali et al. (2022), Avval et al. (2022), Namalomba et al. (2022), Solinas et al. (2021), Jäger (2021), Ibrahim et al. (2021), Du and Xiao (2019), Xu et al. (2019), Harati et al. (2021), Cummings and Crooks (2020), Schauder et al. (2020), Sert et al. (2020), Fuller et al. (2020), Liang et al. (2020), Aghaie and Heidary (2019), Norman et al. (2018), Dehghanpour et al. (2018), Esmaeili Aliabadi et al. (2017), Jalalimanesh, Shahabi Haghighi, et al. (2017), Jalalimanesh, Haghighi, et al. (2017), Dehghanpour et al. (2016), Salle (2015), Lakić et al. (2015)
Validation	Masuda et al. (2022), Jørgensen et al. (2022), Watson et al. (2022), Krivorotko et al. (2022), White et al. (2022), Pawar and Jha (2022), Perumal and Zyl (2022), Kim et al. (2021), Reiker et al. (2021), Ye et al. (2021), Cockrell and An (2021), Zhang et al. (2020), Neri (2019), van der Hoog (2019), Oyebamiji et al. (2019), Boyd et al. (2018), Pearce and Slade (2018), Lamperti et al. (2018), Chu et al. (2018), Fievet and Sornette (2018), Chen et al. (2017), Kattwinkel and Reichert (2017), Chapron et al. (2016), van der Vaart et al. (2016), van der Vaart et al. (2015), Lagarrigues et al. (2015)
Analysis	Marzouk and Hassan (2022), Xie et al. (2022), Ebrie and Kim (2022), Li et al. (2022), Angione et al. (2022), Kieu et al. (2022), Kumar et al. (2021), Shiono (2021), Nunes et al. (2021), Gursoy and Badur (2021), ten Broeke et al. (2021), Xiao and Liu (2021), Chen et al. (2021), Larie et al. (2021), Xu et al. (2020), Koda et al. (2020), Davis et al. (2020), Day et al. (2020), van Strien et al. (2019), Bhattacharjee et al. (2019), Ma et al. (2019), Vahdati et al. (2019), Edali and Yücel (2019), Janssen et al. (2019), Garg et al. (2019), Perry and O'Sullivan (2018), Ozik et al. (2018), Zhang et al. (2018), Scott et al. (2018), MacPherson et al. (2017), Hayashi et al. (2016), Papadopoulos and Azar (2016), Rajabi et al. (2016), Peters et al. (2016)

Abbreviations: ABM, agent-based modelling; ML, machine learning.

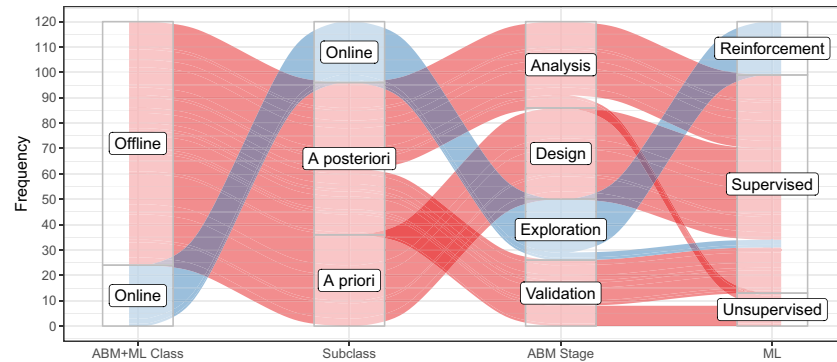


FIGURE 10 Reviewed papers on ML in ABM by application stage and type of ML algorithm. ABM, agent-based modelling; ML, machine learning

TABLE 3 Publications that use ML in ABM to assist the design stage, categorized by the adopted ML algorithm

Algorithm	References
ANN	Aslani et al. (2022), Kadinski et al. (2022), Liang et al. (2022), Peralta et al. (2022), Batata et al. (2019), Candadai et al. (2019), Jäger (2019), Saeedi (2018), Romero-Mujalli et al. (2017), Wolf et al. (2015)
DT	Taghikhah et al. (2022), Salas-Rodríguez and Rivas-Tovar (2022), Rosés et al. (2021), Augustijn et al. (2020), Snyder et al. (2019), Bhattacharjee et al. (2018), Cenek and Franklin (2017), Alonso-Betanzos et al. (2017), Sánchez-Marño et al. (2015), Balasubramaniam and Thigarasu (2015)
RF	Zornić and Marković (2022), Wang et al. (2021), Kavak et al. (2018)
GB	Ghaffarian et al. (2021)
BN	Marvuglia et al. (2022), Abdulkareem et al. (2019), Hyun et al. (2019), Abdulkareem et al. (2018)
LogR	García-Magariño et al. (2020), Lee and Hong (2019), Sadeghipour et al. (2017), Zhang et al. (2016), Lee and Wong (2016)

Abbreviations: ABM, agent-based modelling; ML, machine learning.

ML algorithm, the existence of data from the domain of application is assumed. Applications include the generation of synthetic data that will be used later by the ABM, for example, Saeedi (2018) generates predictions of land use change with neural networks and subsequently implements those results as part of the environment. Other publications with a similar purpose are Ghaffarian et al. (2021); Wang et al. (2021); Yao et al. (2020). A second group of applications has to do with generating the behavioural rules of the agents so that the modeller can later implement them in the model, for example, Augustijn et al. (2020); Alonso-Betanzos et al. (2017); Pouladi et al. (2020). It is important to note that the algorithms adopted for this purpose are usually those that are easy to interpret and allow the rules to be derived in a simple way, for example, decision trees. A third group of applications has to do with learning an ML model from data, and later inserting that model in the ABM so that it is available to agents, for example, Jäger (2019); Romero-Mujalli et al. (2017); Wolf et al. (2015). Unlike in the previous type of application, interpretability is usually discarded for the sake of precision. Adopted ML algorithms include neural networks, random forests, or other similar ensemble methods.

The application domains are mainly concentrated in social sciences (6 publications), for example, modelling social interactions (Candadai et al., 2019; Jäger, 2019; Wolf et al., 2015); medicine (6), for example, modelling the spread of diseases (Abdulkareem et al., 2018; Abdulkareem et al., 2019; Augustijn et al., 2020); and environmental sciences (8), for example, dealing with natural resources management (Hyun et al., 2019; Liang et al., 2022; Marvuglia et al., 2022; Pouladi et al., 2020; Saeedi, 2018).

5.2 | Offline a posteriori learning

Offline a posteriori learning is used for validation and analysis. Supervised ML is mostly used, but unlike the a priori learning case, the training data is generated by the ABM itself. Some exceptions make use of unsupervised learning, for example, Peters et al. (2016) combine an individual-based Plant Interaction Model with self-organizing feature maps (SOM), a type of neural network, to visualize the relationship between inputs and

outputs. Koda et al. (2020) use the Louvain algorithm for analysing the formation of social networks and reconstructing clustering structures. Ye et al. (2021), with the purpose of calibrating parameters to reproduce real behaviour in distributed systems, use PCA to compress the high dimensionality of the agent state space. Marzouk and Hassan (2022) employ the k-means algorithm to cluster the evacuation performance of both museum visitors and workers in response to fire. The supervised ML algorithms used for validation purposes include genetic algorithms (Cockrell & An, 2021; White et al., 2022), Boost DT (Lamperti et al., 2018; Zhang et al., 2020), random forest (Pawar & Jha, 2022), support vector machines (Perumal & Zyl, 2022), simulated annealing (Neri, 2019), neural networks (Jørgensen et al., 2022; Krivorotko et al., 2022; Reiker et al., 2021; van der Hoog, 2019), and Bayesian networks (Chu et al., 2018; Masuda et al., 2022). It is worth noting that the interpretability of the algorithm is not relevant in this case. The goal is to find the subset of parameters that replicate the desired behaviour.

During the analysis stage, as shown in Table 4, there is a clear preference towards neural network algorithms (12 papers), and random forests (9). Applications include understanding the relationship between input and output parameters, for example, Janssen et al. (2019) propose analysing emergence through causal discovery, and Edali and Yücel (2019) use random forests to understand the dynamics of modelled systems. Surrogate models are also of interest, where ML is used to approximate the ABM. For example, ten Broeke et al. (2021) use support vector machines to approximate the ABM simulation in three cases, allowing the distinguishing of which parameters determine the limits in the parameter space between regions; Larie et al. (2021) use an ANN as a surrogate system to predict the time of evolution of an ABM, and Papadopoulos and Azar (2016) propose surrogate models with multiple linear regressions.

Regarding the use of ML for validation in ABM, several studies have used ML techniques to support the validation and calibration of ABMs in the field of economics. For example, Zhang et al. (2020) and Lamperti et al. (2018) both used ML surrogates to reproduce the asset pricing model, while Neri (2019), Chu et al. (2018), and Fievet and Sornette (2018) used ML techniques to approximate financial or economic time series in their ABMs. For the analysis stage, most of the applications are concentrated in medicine (7 papers), biological sciences (6), and social sciences (5). Again, regarding the domain of medicine, the models try to understand the spread of diseases such as COVID-19 (Xiao & Liu, 2021), Cutaneous Leishmaniasis (Rajabi et al., 2016), or sepsis (Larie et al., 2021). For biological sciences, applications try to analyse some aspects of animal societies, for example, groups formation (Koda et al., 2020), mortality and landscape (Day et al., 2020), foraging behaviour (Bhattacharjee et al., 2019), self-thinning rule in plant populations Ma et al. (2019), response to fluctuating food resource availability Scott et al. (2018), reproductive fitness and animal personality MacPherson et al. (2017), and below-ground competition among plants Peters et al. (2016). For social sciences, the problems addressed are related to human behaviour or dynamics, specifically, brain drain (Gursoy & Badur, 2021), late Pleistocene human dispersal in Africa (Vahdati et al., 2019), and the reconstruction of human interactions in past human environments (Perry & O'Sullivan, 2018).

5.3 | Online learning

Online learning is related exclusively to the experimentation stage, allowing agents to explore the environment to generate an optimal behaviour rule. In this case, the data is generated during the simulation. In most cases, some reinforcement learning algorithm has been adopted by the community, although there are some instances in which supervised algorithms have been chosen. For instance, Jäger (2021) introduces a framework for ABM that uses artificial neural networks to learn behaviour rules, and Dehghanpour et al. (2016) propose that each agent, representing a company in the energy market domain, develops a private dynamic Bayesian network of the market using sparse Bayesian learning to train the model that is then used to infer the future state of the market for estimating the optimal bidding function. Arguing that economic models are too simplistic, Salle (2015) introduces an expectation model of inflation based on an artificial neural network. The weights are randomly initialized so that, at least at the beginning, agents' mental models and the resulting expectations differ. As new observations become available, the neural networks

TABLE 4 Publications that use ML in ABM to assist the analysis process by the algorithm used

Algorithm	References
DT	Vahdati et al. (2019), Bhattacharjee et al. (2019), Scott et al. (2018), Hayashi et al. (2016)
RF	Gursoy and Badur (2021), Chen et al. (2021), Day et al. (2020), Ma et al. (2019), Edali and Yücel (2019), Garg et al. (2019), Perry and O'Sullivan (2018), Ozik et al. (2018), MacPherson et al. (2017)
SVM	Ebrie and Kim (2022), ten Broeke et al. (2021), van Strien et al. (2019)
ANN	Angione et al. (2022), Kieu et al. (2022), Kumar et al. (2021), Shiono (2021), Xiao and Liu (2021), Chen et al. (2021), Larie et al. (2021), Davis et al. (2020), Xu et al. (2020), Ozik et al. (2018), Zhang et al. (2018), Peters et al. (2016)
BN	Rajabi et al. (2016)
LinR	Papadopoulos and Azar (2016)
k-means	Marzouk and Hassan (2022), Xie et al. (2022), Li et al. (2022)

Abbreviations: ABM, agent-based modelling; ML, machine learning.

are trained. Q-learning is the most used reinforcement learning, followed by other temporal difference algorithms. Other proposals involve the use of Deep Q-Nets, although the application domain is still in unrealistic environments such as the Schelling Segregation model.

Most of the works in this category adopt the learned optimal behaviour rule as a hypothesis of the studied phenomenon. However, three of them assist the modeller in the design of the ABM. The approach proposed by Cummings and Crooks (2020) produces agents that create their own optimal strategy, which the designer can interpret. Fuller et al. (2020) propose the application of reinforcement learning to enable agents to define their necessary interaction rules, and Norman et al. (2018) examine a novel path to agent policy development in which modellers do not manually craft the policies, but allow them to emerge through the application of reinforcement learning within a game engine environment.

The main application domains of online learning are energy (nine publications) and medicine (3). The electricity markets are characterized by suppliers and consumers of electric energy, and applications vary around the supply strategies and the number of agents that are modelled. For example, Xu et al. (2019) model a residential demand-responsive electricity market in which both power generation companies and retailers can optimize their supply strategy through reinforcement learning. Liang et al. (2020) adopt a deep deterministic policy gradient algorithm to model the bidding strategies of generation companies for scenarios with complete and incomplete information, even in non-discrete spaces. In the medical area, applications are related to finding optimal therapies. For instance, Jalalimanesh et al. (2017,a) use Q-learning to optimize key elements in radiation therapy, such as dose and fractionation schemes, and Schauder et al. (2020) use temporal difference to study habit formation by modelling preference formation in terms of positive or negative feedback about the food most recently consumed.

6 | DISCUSSION

To discuss future trends in the application of ML in ABM, we have examined the progression of stages where ML is utilized and the evolution of ML techniques over time. Figure 11 summarizes the number of annual publications according to the ABM stage involved. The trend shows an increasing number of proposals, from seven publications in 2015 to 26 in 2022. These results confirm the relevance of studying recent years to get a complete picture of the state of the art in the field, which is in line with the exploratory analysis shown in Figure 6. The design and validation stages were the most supported by ML in the last year. Offline *a posteriori* learning has accounted for between 50% and 60% of the publications in the last 2 years, denoting an interest in the validation and analysis stages. In contrast, the interest in online learning has gradually decreased from 35.7% in 2020 to 14.4% in the last year.

Figure 11 also provides an in-depth study of the trends in the adoption of ML algorithms. ANNs are the most used supervised learning algorithm, with 29 proposals. The number of articles that used some ANN algorithm increased by more than 100% during the period 2020–2021 compared to the period 2015–2019. Interestingly, these learning models, considered black boxes with poor interpretability, are also being used in the analysis stage. Metaheuristic-based learning algorithms, such as genetic algorithms, are being adopted more frequently, primarily for parameter

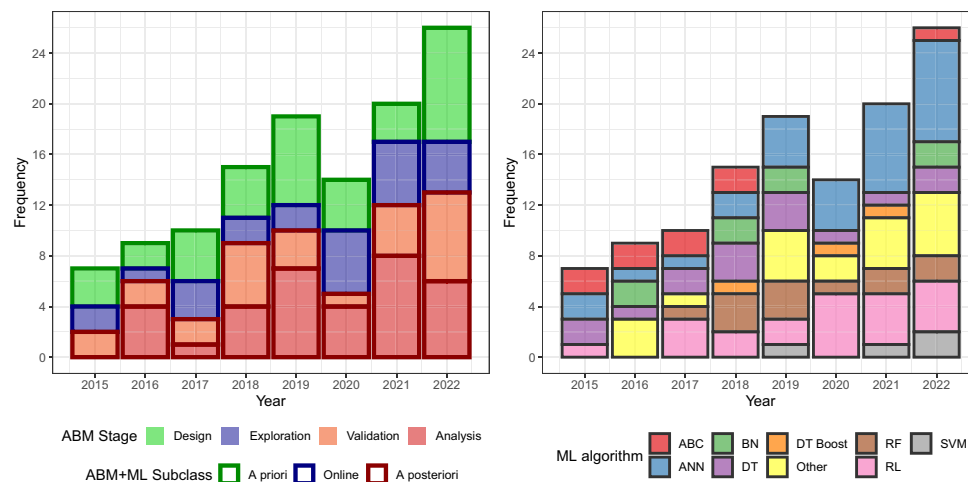


FIGURE 11 Analysis of the number of proposals reviewed. Proposals per year and ABM stages (left); Proposals per year and ML algorithm (right). ABM, agent-based modelling; ML, machine learning

validation or calibration tasks. Similarly, unsupervised learning algorithms, such as k-means and PCA, have started to be utilized since 2020, typically for analysing simulation results. In contrast, the adoption of algorithms with better interpretability, such as decision trees, Bayesian networks, and logistic regression, has decreased. Finally, the number of proposals that adopt reinforcement learning has also increased, with Q-learning being the preferred algorithm. However, some proposals with other types of reinforcement algorithms that are capable of contending with more complex environments are currently being explored.

The use of ML in ABM modelling methodology will continue to increase in the coming years, due to the growing availability of data and computing power, allowing the development of more complex models that can be validated using ML techniques. Furthermore, the increasing adoption of ML in other fields is likely to inspire more researchers to apply these techniques to ABM as well.

Regarding new challenges, various pertinent but unexplored ML algorithms still need exploration, for example, unsupervised learning algorithms have been little explored in any of the ABM stages other than validation, either for clustering or dimensional reduction. Reinforcement learning dominates the explorations for online learning, despite the pertinence of incremental supervised and unsupervised algorithms in this regard. Examples of the latter include online clustering, online representation learning, and online anomaly detection tasks. A comprehensive review of online learning algorithms can be found in the work of Hoi et al. (2021). Online clustering algorithms could detect emergent patterns in simulation results, while online representation learning algorithms could extract meaningful features from high-dimensional data. Online anomaly detection algorithms could monitor the behaviour of agents in real-time and detect abnormal behaviour that could indicate a problem with the model. The potential benefits of these algorithms include increased efficiency and accuracy in model validation and analysis. Another unexplored area is the use of deep learning algorithms such as convolutional neural networks in ABM. These algorithms have shown great success in image and signal processing tasks and could potentially extract meaningful features from agent-based simulation data as well.

There are several novel ML approaches that could be useful in ABM modelling. One is active learning, which improves the efficiency of the learning process by selecting specific instances from the dataset to be labelled. This approach could reduce the labelling cost in ABM model validation, where a large volume of labelled data is often required. Another approach that could be useful is transfer learning, which allows the use of previously acquired knowledge in one domain to improve learning in another related domain. This approach could be applied in ABM modelling to transfer knowledge from one model to another or to use previously acquired information in constructing similar models. Other novel ML approaches that could be applied in ABM modelling include imitation learning and self-supervised learning. These approaches could address different challenges in ABM model construction, such as identifying emergent patterns and adapting to changing environments. Overall, applying these novel ML approaches in ABM modelling could improve the efficiency and accuracy of model construction, allowing the identification of complex patterns and trends in the systems under study. However, it is important to note that these approaches may require greater knowledge and expertise in model design and implementation, as well as data processing and analysis.

Another challenge, even more critical, is related to the poor documentation of implementations. In the worst case, the description of both the ML algorithm used and the ABM model is minimal. In the best cases, an attempt was made to describe both tools separately. However, to assist less experienced developers in the area, there is a lack of a unified description that allows the integration of both tools to be understood as one. This would enable the design of a clear and transmittable workflow for each of the problems to be addressed.

We propose extending the ODD protocol (Overview, Design concepts, Details; Grimm et al., 2010, 2020) to include a standardized description of the integration of ML and ABM. ODD is a standardized document that provides a consistent, logical, and readable account of the structure and dynamics of ABMs. ODD consists of seven elements as shown in Figure 12. However, the current ODD Protocol may not be able to accommodate the diversity of applications of ML in ABMs. Fortunately, the ODD Protocol covers other necessary elements to carry out such descriptions, making its extension for documenting the diverse applications of ML identified in the ABM literature simpler. We propose categorizing ML applications in ABMs to enable the design questions guiding the extension to the learning component of the protocol. For instance, we suggest including the following questions: What is the purpose of learning? When is learning performed? What components of the ABM are affected by learning? How is learning computed?

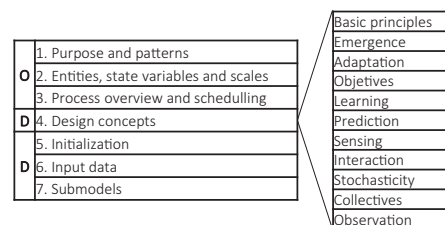


FIGURE 12 Structure of model descriptions following the Overview, Design concepts, Details protocol



Additionally, other ABM documentation methods, such as STRESS (Monks et al., 2018) and RAT-RS (Achter et al., 2022), may also provide valuable insights for the extension of the ODD Protocol. Therefore, we suggest exploring the potential contributions of these methods to ensure that the extended protocol is comprehensive and covers all relevant aspects of integrating ML and ABM.

7 | CONCLUSIONS

The use of ML methods in ABM will become increasingly important since scientists intend to adopt more realistic models with higher complexity and computational requirements. Inherently, the difficulties for designing, calibrating, and analysing such models also increase, suggesting the adoption of ML methods to assist these tasks. This work presents a novel review of the latest trends in ABM assisted by ML. The review includes 120 articles published in high-impact journals from 2015 to 2022. The works were summarized based on the problems that could be solved through the application of ML in ABM. The novelty of this review lies in the fact that each of these problems or criticisms was related to the stages of the ABM development cycle, allowing modellers with less experience in this approach to deduce an adequate workflow when carrying out their implementations. For this, we propose the following categorization:

Offline learning: In its two forms:

1. *A priori:* Before the simulation is carried out, the learning algorithm is applied to existing data of the phenomenon to be modelled, resolving to a lesser or greater extent aspects related to the model's design. There are mainly three types of applications: instantiating data in the model, deriving behaviour rules to be implemented in the model, and encapsulating the simulation within the agents' available knowledge.
2. *A posteriori:* After the simulation is completed, ML algorithms are applied to the model's outputs. This approach can address two of the criticisms of ABMs: validation/calibration problems and analysis problems. The former refers to finding a subset of input parameters that generate the desired output, while the latter aims to explain the relationship between inputs and outputs.

Online learning: During the simulation, learning is applied to the data generated by the agents as they interact with previously defined reward structures and environmental states to find optimal behaviour rules.

The use of ML methods in ABM offers several advantages. First, it can help reduce the computational cost associated with simulating complex models, allowing modellers to obtain more realistic and accurate results in a shorter amount of time. Second, ML methods can aid in the design and calibration of the model, which are often difficult and time-consuming tasks. Third, ML can help better analyse and understand the relationships between inputs and outputs, as well as identify patterns and behaviours that may not be immediately obvious.

Despite the potential benefits of using ABM supported by ML, a significant problem remains in the lack of understanding of ML methods. This lack of knowledge can lead to problems and omissions in reporting results, ultimately hindering the advancement of ABMs as research tools. While the implementation of ML algorithms is a complex process involving various steps, some researchers fail to provide sufficient detail about the steps taken to choose a particular model. This lack of transparency can lead to ambiguity, which undermines the credibility of the research.

Additionally, the use of certain ML techniques, such as neural networks, can be driven more by their popularity than their intrinsic utility for solving a particular problem. This trend highlights the need for researchers to be more critical and selective in their use of ML techniques, ensuring they choose the most appropriate technique for their specific research question, rather than simply following the latest trend. By doing so, researchers can improve the quality and accuracy of their results and contribute to the advancement of the field.

One significant criticism of existing approaches for integrating ABM and ML is their inability to propose general methods that apply to many problems and domains. The lack of a clear and systematic methodological framework for implementing these systems compromises their adaptability to face new problems. This limitation highlights the need for a workflow that allows for more transparent and easily replicable models, which can provide greater certainty about the advantages of adopting the ABM approach supported by ML.

Developing such a workflow presents a significant challenge as it requires the integration of multiple components, including data acquisition, pre-processing, model selection, validation, and interpretation. Furthermore, it is essential to consider the specific characteristics of the problem domain, including available data and existing knowledge. The workflow should be transparent and easily replicable to enable other researchers to build on previous work and compare results across studies. Another significant challenge is ensuring the workflow is flexible and adaptable to different problem domains and datasets. This adaptability will require the development of modular components that can be easily swapped in and out, depending on the specific needs of the problem at hand.

AUTHOR CONTRIBUTIONS

Alejandro Platas-López: Writing-original draft; writing-review; visualization and editing. **Alejandro Guerra-Hernández:** Writing-original draft; writing-review and editing. **Marcela Quiroz-castellanos:** Writing-review and editing. **Nicandro Cruz-Ramírez:** Writing-review and editing.

ACKNOWLEDGEMENT

The first author acknowledges support from CONACyT through a scholarship No. 743662.

CONFLICT OF INTEREST STATEMENT

The authors declare no conflict of interest.

DATA AVAILABILITY STATEMENT

Data sharing is not applicable to this article as no new data were created or analyzed in this study.

ORCID

Alejandro Platas-López  <https://orcid.org/0000-0002-2584-343X>

REFERENCES

- Abdi, H., & Williams, L. J. (2010). Principal component analysis. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2(4), 433–459.
- Abdulkareem, S., Augustijn, E.-W., Mustafa, Y., & Filatova, T. (2018). Intelligent judgements over health risks in a spatial agent-based model. *International Journal of Health Geographics*, 17(1), 8.
- Abdulkareem, S., Mustafa, Y., Augustijn, E.-W., & Filatova, T. (2019). Bayesian networks for spatial learning: A workflow on using limited survey data for intelligent learning in spatial agent-based models. *Geoinformatica*, 23(2), 243–268.
- Achter, S., Borit, M., Chattoe-Brown, E., & Siebers, P.-O. (2022). RAT-RS: A reporting standard for improving the documentation of data use in agent-based modelling. *International Journal of Social Research Methodology*, 25(4), 517–540.
- Aghaie, A., & Heidary, M. (2019). Simulation-based optimization of a stochastic supply chain considering supplier disruption: Agent-based modeling and reinforcement learning. *Scientia Iranica*, 26(6), 3780–3795.
- Ali, G. G., El-adaway, I. H., Sims, C., Holladay, J. S., & Chen, C.-F. (2022). Studying dynamic pricing in electrical power markets with distributed generation: Agent-based modeling and reinforcement-learning approach. *Journal of Energy Engineering*, 148(5), 1–22.
- Alonso-Betanzos, A., Sánchez-Maróño, N., Fontenla-Romero, O., Polhill, J. G., Craig, T., Bajo, J., & Corchado, J. M. (Eds.). (2017). *Agent-based modeling of sustainable behaviors* (pp. 53–76). Springer International Publishing.
- Angione, C., Silverman, E., & Yaneske, E. (2022). Using machine learning as a surrogate model for agent-based simulations. *PLoS One*, 17(2), e0263150.
- Ashreeta, P., Holzhauer, S., & Krebs, F. (2019). Overview of machine learning and data-driven methods in agent-based modeling of energy markets. In K. David & K. Geihs (Eds.), *INFORMATIK 2019: 50 Jahre Gesellschaft für Informatik – Informatik für Gesellschaft* (Vol. 294, pp. 571–584). Gesellschaft für Informatik (GI).
- Aslani, Z. H., Omidvar, B., & Karbassi, A. (2022). Integrated model for land-use transformation analysis based on multi-layer perception neural network and agent-based model. *Environmental Science and Pollution Research*, 29(39), 59770–59783.
- Augustijn, E.-W., Abdulkareem, S., Sadiq, M., & Albabawat, A. (2020). Machine learning to derive complex behaviour in agent-based modelling. In *In 2020 International Conference on Computer Science and Software Engineering (CSASE)* (pp. 284–289). Institute of Electrical and Electronics Engineers.
- Avval, A. E., Dehghanian, F., & Pirayesh, M. (2022). Auction design for the allocation of carbon emission allowances to supply chains via multiagent-based model and q-learning. *Computational and Applied Mathematics*, 41(4), 170–175.
- Axtell, R. L. (2000). Why agents? On the varied motivations for agent computing in the social sciences. In *Workshop on agent simulation: Applications, models, and tools*. University of Chicago.
- Balasubramaniam, C., & Thigarasu, V. (2015). Agent-based modeling in supply chain management using improved c4-5. *Research Journal of Applied Sciences, Engineering and Technology*, 9(2), 91–97.
- Baldwin, W. C., Sauser, B., & Cloutier, R. (2015). Simulation approaches for system of systems: Events-based versus agent based modeling. *Procedia Computer Science*, 44, 363–372.
- Batata, O., Augusto, V., & Xie, X. (2019). Mixed machine learning and agent-based simulation for respite care evaluation. In *In 2018 Winter Simulation Conference (WSC)* (pp. 2668–2679). Institute of Electrical and Electronics Engineers Inc..
- Bhattacharjee, S., MacPherson, B., & Gras, R. (2018). A comparison of sexual selection versus random selection with respect to extinction and speciation rates using individual based modeling and machine learning. *Ecological Complexity*, 36, 126–137.
- Bhattacharjee, S., MacPherson, B., Wang, R. F., & Gras, R. (2019). Animal communication of fear and safety related to foraging behavior and fitness: An individual-based modeling approach. *Ecological Informatics*, 54, 101011.
- Boyd, R., Roy, S., Sibily, R., Thorpe, R., & Hyder, K. (2018). A general approach to incorporating spatial and temporal variation in individual-based models of fish populations with application to Atlantic mackerel. *Ecological Modelling*, 382, 9–17.
- Breckling, B. (2002). Individual-based modelling potentials and limitations. *The Scientific World Journal*, 2, 1044–1062.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
- Broniec, W., An, S., Rugarber, S., & Goel, A. K. (2021). Guiding parameter estimation of agent-based modeling through knowledge-based function approximation. In *Spring Symposium on Combining Machine Learning and Knowledge Engineering*. Central Europe (CEUR).
- Candadaí, M., Setzler, M., Izquierdo, E., & Froese, T. (2019). Embodied dyadic interaction increases complexity of neural dynamics: A minimal agent-based simulation model. *Frontiers in Psychology*, 10, 540–549.
- Catullo, E., Gallegati, M., & Russo, A. (2022). Forecasting in a complex environment: Machine learning sales expectations in a stock flow consistent agent-based simulation model. *Journal of Economic Dynamics and Control*, 139, 104405.
- Cenek, M., & Franklin, M. (2017). An adaptable agent-based model for guiding multi-species pacific salmon fisheries management within a SES framework. *Ecological Modelling*, 360, 132–149.
- Chapron, G., Wikenros, C., Liberg, O., Wabakken, P., Flagstad, Ø., Milleret, C., Månsson, J., Svensson, L., Zimmermann, B., Åkesson, M., & Sand, H. (2016). Estimating wolf (*canis lupus*) population size from number of packs and an individual based model. *Ecological Modelling*, 339, 33–44.



- Charteris, P., Golden, B., & Garrick, D. J. (2001). Livestock breeding industries as complex adaptive systems. *Conference of the Association for the Advancement of Animal Breeding and Genetics*, 14, 461–464.
- Chen, C.-M., Drovandi, C., Keith, J., Anthony, K., Caley, M., & Mengersen, K. (2017). Bayesian semi-individual based model with approximate Bayesian computation for parameters calibration: Modelling crown-of-thorns populations on the great barrier reef. *Ecological Modelling*, 364, 113–123.
- Chen, S., Londoño-Larrea, P., McGough, A., Bible, A., Gunaratne, C., Araujo-Granda, P., Morrell-Falvey, J., Bhowmik, D., & Fuentes-Cabrera, M. (2021). Application of machine learning techniques to an agent-based model of pantoea. *Frontiers in Microbiology*, 12, 1–10.
- Chu, Z., Yang, B., Ha, C., & Ahn, K. (2018). Modeling gdp fluctuations with agent-based model. *Physica A: Statistical Mechanics and its Applications*, 503, 572–581.
- Cockrell, C., & An, G. (2021). Utilizing the heterogeneity of clinical data for model refinement and rule discovery through the application of genetic algorithms to calibrate a high-dimensional agent-based model of systemic inflammation. *Frontiers in Physiology*, 12, 1–12.
- Conte, R., & Paolucci, M. (2014). On agent-based modeling and computational social science. *Frontiers in Psychology*, 5, 668.
- Cummings, P., & Crooks, A. (2020). Development of a hybrid machine learning agent based model for optimization and interpretability. *Lecture Notes in Computer Science*, 12268, 151–160.
- Davis, C. N., Hollingsworth, T. D., Caudron, Q., & Irvine, M. A. (2020). The use of mixture density networks in the emulation of complex epidemiological individual-based models. *PLoS Computational Biology*, 16(3), e1006869.
- Day, C. C., Zollner, P. A., Gilbert, J. H., & McCann, N. P. (2020). Individual-based modeling highlights the importance of mortality and landscape structure in measures of functional connectivity. *Landscape Ecology*, 35(10), 2191–2208.
- Dehghanpour, K., Hashem Nehrir, M., Sheppard, J., & Kelly, N. (2018). Agent-based modeling of retail electrical energy markets with demand response. *IEEE Transactions on Smart Grid*, 9(4), 3465–3475.
- Dehghanpour, K., Nehrir, M., Sheppard, J., & Kelly, N. (2016). Agent-based modeling in electrical energy markets using dynamic bayesian networks. *IEEE Transactions on Power Systems*, 31(6), 4744–4754.
- Du, H., & Xiao, T. (2019). Pricing strategies for competing adaptive retailers facing complex consumer behavior: Agent-based model. *International Journal of Information Technology & Decision Making*, 18(6), 1909–1939.
- Ebrie, A. S., & Kim, Y. J. (2022). Investigating market diffusion of electric vehicles with experimental design of agent-based modeling simulation. *Systems*, 10(2), 28.
- Edali, M., & Yücel, G. (2019). Exploring the behavior space of agent-based simulation models using random forest metamodells and sequential sampling. *Simulation Modelling Practice and Theory*, 92, 62–81.
- Efron, S., & Ravid, R. (2018). *Writing the literature review: A practical guide*. Guilford Publications.
- Esmaili Aliabadi, D., Kaya, M., & Sahin, G. (2017). Competition, risk and learning in electricity markets: An agent-based simulation study. *Applied Energy*, 195, 1000–1011.
- Fievet, L., & Sornette, D. (2018). Calibrating emergent phenomena in stock markets with agent based models. *PLoS One*, 13(3), e0193290.
- Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *The Annals of Statistics*, 29(5), 1189–1232.
- Friedman, N., Geiger, D., & Goldszmidt, M. (1997). Bayesian network classifiers. *Machine Learning*, 29(2/3), 131–163.
- Fuller, D., de Arruda, E., & Ferreira Filho, V. (2020). Learning-agent-based simulation for queue network systems. *Journal of the Operational Research Society*, 71(11), 1723–1739.
- García-Magariño, I., Medrano, C., & Delgado, J. (2020). Estimation of missing prices in real-estate market agent-based simulations with machine learning and dimensionality reduction methods. *Neural Computing and Applications*, 32(7), 2665–2682.
- Garg, A., Yuen, S., Seekhao, N., Yu, G., Karwowski, J., Powell, M., Sakata, J., Mongeau, L., Jaja, J., & Li-Jessen, N. (2019). Towards a physiological scale of vocal fold agent-based models of surgical injury and repair: Sensitivity analysis, calibration and verification. *Applied Sciences (Switzerland)*, 9(15), 1–32.
- Ghaffarian, S., Roy, D., Filatova, T., & Kerle, N. (2021). Agent-based modelling of post-disaster recovery with remote sensing data. *International Journal of Disaster Risk Reduction*, 60, 102285.
- Ghorbani, A., Dechesne, F., Dignum, V., & Jonker, C. (2014). Enhancing ABM into an inevitable tool for policy analysis. *Journal on Policy and Complex Systems*, 1(1), 61–76.
- Gilbert, N. (2020). *Agent-based models*. SAGE Publications.
- Goodell, J. W., Kumar, S., Lim, W. M., & Pattnaik, D. (2021). Artificial intelligence and machine learning in finance: Identifying foundations, themes, and research clusters from bibliometric analysis. *Journal of Behavioral and Experimental Finance*, 32, 100577.
- Grimm, V. (1999). Ten years of individual-based modelling in ecology: What have we learned and what could we learn in the future? *Ecological Modelling*, 115(2–3), 129–148.
- Grimm, V., Berger, U., DeAngelis, D. L., Polhill, J. G., Giske, J., & Railsback, S. F. (2010). The odd protocol: A review and first update. *Ecological Modelling*, 221(23), 2760–2768.
- Grimm, V., Railsback, S. F., Vincenot, C. E., Berger, U., Gallagher, C., DeAngelis, D. L., Edmonds, B., Ge, J., Giske, J., Groeneveld, J., Johnston, A. S., Milles, A., Nabe-Nielsen, J., Polhill, G., Radchuk, V., Rohwäder, M.-S., Stillman, R. A., Thiele, J. C., & Ayllón, D. (2020). The odd protocol for describing agent-based and other simulation models: A second update to improve clarity, replication, and structural realism. *Journal of Artificial Societies and Social Simulation*, 23(2), 7.
- Gursoy, F., & Badur, B. (2021). An agent-based modeling approach to brain drain. *IEEE Transactions on Computational Social Systems*, 9(2), 356–365.
- Harati, S., Perez, L., & Molowny-Horas, R. (2021). Promoting the emergence of behavior norms in a principal-agent problem—An agent-based modeling approach using reinforcement learning. *Applied Sciences (Switzerland)*, 11(18), 1–24.
- Hayashi, S., Prasasti, N., Kanamori, K., & Ohwada, H. (2016). Improving behavior prediction accuracy by using machine learning for agent-based simulation. *Lecture Notes in Computer Science*, 9621, 280–289.
- Haykin, S., Haykin, S., & Haykin, S. (1999). *Neural networks: A comprehensive foundation*. Prentice Hall.
- Hearst, M., Dumais, S., Osuna, E., Platt, J., & Scholkopf, B. (1998). Support vector machines. *IEEE Intelligent Systems and their Applications*, 13(4), 18–28.
- Heppenstall, A., Crooks, A., Malleson, N., Manley, E., Ge, J., & Batty, M. (2020). Future developments in geographical agent-based models: Challenges and opportunities. *Geographical Analysis*, 53(1), 76–91.
- Hoi, S. C., Sahoo, D., Lu, J., & Zhao, P. (2021). Online learning: A comprehensive survey. *Neurocomputing*, 459, 249–289.

- Hyun, J.-Y., Huang, S.-Y., Yang, Y.-C., Tidwell, V., & Macknick, J. (2019). Using a coupled agent-based modeling approach to analyze the role of risk perception in water management decisions. *Hydrology and Earth System Sciences*, 23(5), 2261–2278.
- Ibrahim, M., Hashmi, U., Nabeel, M., Imran, A., & Ekin, S. (2021). Embracing complexity: Agent-based modeling for hetnets design and optimization via concurrent reinforcement learning algorithms. *IEEE Transactions on Network and Service Management*, 18, 4042–4062.
- Jäger, G. (2019). Replacing rules by neural networks a framework for agent-based modelling. *Big Data and Cognitive Computing*, 3(4), 1–12.
- Jäger, G. (2021). Using neural networks for a universal framework for agent-based models. *Mathematical and Computer Modelling of Dynamical Systems*, 27(1), 162–178.
- Jalalimanesh, A., Haghighi, H., Ahmadi, A., Hejazian, H., & Soltani, M. (2017). Multi-objective optimization of radiotherapy: Distributed q-learning and agent-based simulation. *Journal of Experimental and Theoretical Artificial Intelligence*, 29(5), 1071–1086.
- Jalalimanesh, A., Shahabi Haghighi, H., Ahmadi, A., & Soltani, M. (2017). Simulation-based optimization of radiotherapy: Agent-based modeling and reinforcement learning. *Mathematics and Computers in Simulation*, 133, 235–248.
- Janssen, S., Sharpanskykh, A., Curran, R., & Langendoen, K. (2019). Using causal discovery to analyze emergence in agent-based models. *Simulation Modelling Practice and Theory*, 96, 101940.
- Jørgensen, A. C. S., Ghosh, A., Sturrock, M., & Shahrezaei, V. (2022). Efficient bayesian inference for stochastic agent-based models. *PLoS Computational Biology*, 18(10), e1009508.
- Judson, O. P. (1994). The rise of the individual-based model in ecology. *Trends in Ecology and Evolution*, 9(1), 9–14.
- Kadinski, L., Salcedo, C., Boccelli, D. L., Berglund, E., & Ostfeld, A. (2022). A hybrid data-driven-agent-based modelling framework for water distribution systems contamination response during COVID-19. *Water*, 14(7), 1088.
- Kattwinkel, M., & Reichert, P. (2017). Bayesian parameter inference for individual-based models using a particle markov chain Monte Carlo method. *Environmental Modelling & Software*, 87, 110–119.
- Kavak, H., Padilla, J., Lynch, C., & Dially, S. (2018). Big data, agents, and machine learning: Towards a data-driven agent-based modeling approach. In E. Frydenlund & S. Shafer (Eds.), *ANSS 201818: Proceedings of the annual simulation symposium* (Vol. 50, pp. 125–136). The Society for Modeling and Simulation International.
- Kieu, M., Nguyen, H., Ward, J. A., & Malleon, N. (2022). Towards real-time predictions using emulators of agent-based models. *Journal of Simulation*, 1–18.
- Kim, D., Yun, T.-S., Moon, I.-C., & Bae, J. (2021). Automatic calibration of dynamic and heterogeneous parameters in agent-based models. *Autonomous Agents and Multi-Agent Systems*, 35(2), 1–66.
- Koda, H., Arai, Z., & Matsuda, I. (2020). Agent-based simulation for reconstructing social structure by observing collective movements with special reference to single-file movement. *PLoS ONE*, 15, 1–16.
- Kostadinov, F., Holm, S., Steubing, B., Thees, O., & Lemm, R. (2014). Simulation of a swiss wood fuel and roundwood market: An explorative study in agent-based modeling. *Forest Policy and Economics*, 38, 105–118.
- Krivorotko, O., Sosnovskaia, M., Vashchenko, I., Kerr, C., & Lesnic, D. (2022). Agent-based modeling of COVID-19 outbreaks for New York state and UK: Parameter identification algorithm. *Infectious Disease Modelling*, 7(1), 30–44.
- Kumar, S. A., García-Magariño, I., Nasralla, M. M., & Nazir, S. (2021). Agent-based simulators for empowering patients in self-care programs using mobile agents with machine learning. *Mobile Information Systems*, 2021, 1–10.
- Lagarigues, G., Jabot, F., Lafond, V., & Courbaud, B. (2015). Approximate bayesian computation to recalibrate individual-based models with population data: Illustration with a forest simulation model. *Ecological Modelling*, 306, 278–286.
- Lakić, G., Artač, G., & Gubina, A. (2015). Agent-based modeling of the demand-side system reserve provision. *Electric Power Systems Research*, 124, 85–91.
- Lamperti, F., Roventini, A., & Sani, A. (2018). Agent-based model calibration using machine learning surrogates. *Journal of Economic Dynamics and Control*, 90, 366–389.
- Larie, D., An, G., & Cockrell, R. (2021). The use of artificial neural networks to forecast the behavior of agent-based models of pathophysiology: An example utilizing an agent-based model of sepsis. *Frontiers in Physiology*, 12, 1–10.
- Lee, J.-S., Filatova, T., Ligmann-Zielinska, A., Hassani-Mahmoei, B., Stonedahl, F., Lorscheid, I., Voinov, A., Polhill, J. G., Sun, Z., & Parker, D. C. (2015). The complexities of agent-based modeling output analysis. *Journal of Artificial Societies and Social Simulation*, 18(4), 4.
- Lee, M., & Hong, T. (2019). Hybrid agent-based modeling of rooftop solar photovoltaic adoption by integrating the geographic information system and data mining technique. *Energy Conversion and Management*, 183, 266–279.
- Lee, T.-C., & Wong, K. (2016). An agent-based model for queue formation of powered two-wheelers in heterogeneous traffic. *Physica A: Statistical Mechanics and its Applications*, 461, 199–216.
- Lengnick, M. (2013). Agent-based macroeconomics: A baseline model. *Journal of Economic Behavior & Organization*, 86, 102–120.
- Li, L., Wang, J., Zhong, X., Lin, J., Wu, N., Zhang, Z., Meng, C., Wang, X., Shah, N., Brandon, N., Xie, S., & Zhao, Y. (2022). Combined multi-objective optimization and agent-based modeling for a 100% renewable island energy system considering power-to-gas technology and extreme weather conditions. *Applied Energy*, 308, 118376.
- Li, X., Mao, W., Zeng, D., & Wang, F.-Y. (2008). Agent-based social simulation and modeling in social computing. In *Intelligence and security informatics* (pp. 401–412). Springer.
- Liang, X., Lu, T., & Yishake, G. (2022). How to promote residents' use of green space: An empirically grounded agent-based modeling approach. *Urban Forestry & Urban Greening*, 67, 127435.
- Liang, Y., Guo, C., Ding, Z., & Hua, H. (2020). Agent-based modeling in electricity market using deep deterministic policy gradient algorithm. *IEEE Transactions on Power Systems*, 35(6), 4180–4192.
- Ligmann, A., & Sun, L. (2010). Applying time-dependent variance-based global sensitivity analysis to represent the dynamics of an agent-based model of land use change. *International Journal of Geographical Information Science*, 24, 1829–1850.
- Ma, P., Han, X.-H., Lin, Y., Moore, J., Guo, Y.-X., & Yue, M. (2019). Exploring the relative importance of biotic and abiotic factors that alter the self-thinning rule: Insights from individual-based modelling and machine-learning. *Ecological Modelling*, 397, 16–24.
- MacPherson, B., Mashayekhi, M., Gras, R., & Scott, R. (2017). Exploring the connection between emergent animal personality and fitness using a novel individual-based model and decision tree approach. *Ecological Informatics*, 40, 81–92.
- Malleon, N., Heppenstall, A., & See, L. (2010). Crime reduction through simulation: An agent-based model of burglary. *Computers, Environment and Urban Systems*, 34(3), 236–250.



- Marvuglia, A., Bayram, A., Baustert, P., Gutiérrez, T. N., & Igos, E. (2022). Agent-based modelling to simulate farmers' sustainable decisions: Farmers' interaction and resulting green consciousness evolution. *Journal of Cleaner Production*, 332, 129847.
- Marzouk, M., & Hassan, F. (2022). Modeling evacuation and visitation proximity in museums using agent-based simulation. *Journal of Building Engineering*, 56, 104794.
- Masuda, S., Bahr, K., Tsuchiya, N., & Takemori, T. (2022). Agent based simulation with data driven parameterization for evaluation of social acceptance of a geothermal development: A case study in tsuchiya, Fukushima, Japan. *Scientific Reports*, 12(1), 3314.
- Miller, M. Z., Griendling, K., & Mavris, D. N. (2012). Exploring human factors effects in the smart grid system of systems demand response. In 2012 7th international conference on system of systems engineering (SoSE) (pp. 1–6). IEEE.
- Monks, T., Currie, C. S. M., Onggo, B. S., Robinson, S., Kunc, M., & Taylor, S. J. E. (2018). Strengthening the reporting of empirical simulation studies: Introducing the STRESS guidelines. *Journal of Simulation*, 13(1), 55–67.
- Montgomery, D., Peck, E., & Vining, G. (2015). Introduction to linear regression analysis. In *Wiley series in probability and statistics*. Wiley.
- Namalomba, E., Feihu, H., & Shi, H. (2022). Agent based simulation of centralized electricity transaction market using bi-level and q-learning algorithm approach. *International Journal of Electrical Power & Energy Systems*, 134, 107415.
- Navarro, J. F., Frenk, C. S., & White, S. D. M. (1997). A universal density profile from hierarchical clustering. *The Astrophysical Journal*, 490(2), 493–508.
- Neri, F. (2019). Combining machine learning and agent based modeling for gold price prediction. *Communications in Computer and Information Science*, 900, 91–100.
- Norman, M., Koehler, M., Kutarnia, J., Silvey, P., Tolk, A., & Tracy, B. (2018). Applying complexity science with machine learning, agent-based models, and game engines: Towards embodied complex systems engineering. In *Unifying themes in complex systems IX* (pp. 173–183). Springer.
- Nourqolipour, R., & Shariff, R. (2010). How agent based modeling (abm) can be linked to gis for modelling land use and land cover change. In *MRSS 6th International Remote Sensing and GIS Conference and Exhibition*.
- Nunes, A., Zwick, M., & Wakeland, W. (2021). Sensitivity analysis of an agent-based simulation model using reconstructability analysis. *International Journal of General Systems*, 50(3), 319–338.
- Oyebamiji, O., Wilkinson, D., Li, B., Jayathilake, P., Zuliani, P., & Curtis, T. (2019). Bayesian emulation and calibration of an individual-based model of microbial communities. *Journal of Computational Science*, 30, 194–208.
- Ozik, J., Collier, N., Wozniak, J., Macal, C., & An, G. (2018). Extreme-scale dynamic exploration of a distributed agent-based model with the emews framework. *IEEE Transactions on Computational Social Systems*, 5(3), 884–895.
- Papadopoulos, S., & Azar, E. (2016). Integrating building performance simulation in agent-based modeling using regression surrogate models: A novel human-in-the-loop energy modeling approach. *Energy and Buildings*, 128, 214–223.
- Park, D., & Ryu, D. (2022). Supply chain ethics and transparency: An agent-based model approach with q-learning agents. *Managerial and Decision Economics*, 43(8), 3331–3337.
- Pawar, S., & Jha, A. K. (2022). Analysis of residential location choices of different socio-economic groups and their impact on the density in a city using agent based modelling. *Applied Spatial Analysis and Policy*, 16(1), 119–139.
- Pearce, P., & Slade, R. (2018). Feed-in tariffs for solar microgeneration: Policy evaluation and capacity projections using a realistic agent-based model. *Energy Policy*, 116, 95–111.
- Peralta, A. A., Balta-Ozkan, N., & Longhurst, P. (2022). Spatio-temporal modelling of solar photovoltaic adoption: An integrated neural networks and agent-based modelling approach. *Applied Energy*, 305, 117949.
- Perry, G., & O'Sullivan, D. (2018). Identifying narrative descriptions in agent-based models representing past human-environment interactions. *Journal of Archaeological Method and Theory*, 25(3), 795–817.
- Perumal, R., & Zyl, T. (2022). Surrogate-assisted strategies: The parameterisation of an infectious disease agent-based model. *Neural Computing and Applications*, 34, 1–11.
- Peters, R., Lin, Y., & Berger, U. (2016). Machine learning meets individual-based modelling: Self-organising feature maps for the analysis of below-ground competition among plants. *Ecological Modelling*, 326, 142–151.
- Platas-López, A., Guerra-Hernández, A., Quiroz-Castellanos, M., & Cruz-Ramírez, N. (2023). Agent-based models assisted by supervised learning: A proposal for model specification. *Electronics*, 12(3), 495.
- Pouladi, P., Afshar, A., Molajou, A., & Afshar, M. (2020). Socio-hydrological framework for investigating farmers' activities affecting the shrinkage of urmia lake: hybrid data mining and agent-based modelling. *Hydrological Sciences Journal*, 65(8), 1249–1261.
- Pritchard, A. (1969). Statistical bibliography or bibliometrics. *Journal of Documentation*, 25(4), 348–349.
- Quinlan, J. R. (1986). Induction of decision trees. *Machine Learning*, 1(1), 81–106.
- Vahdati, A., Weissmann, J., Timmermann, A., Ponce de León, M., & Zollikofer, C. (2019). Drivers of late pleistocene human survival and dispersal: An agent-based modeling and machine learning approach. *Quaternary Science Reviews*, 221, 105867.
- Railsback, S., & Grimm, V. (2019). *Agent-based and individual-based modeling: A practical introduction*. Princeton University Press.
- Rajabi, M., Pilešj, P., Shirzadi, M., Fadaei, R., & Mansourian, A. (2016). A spatially explicit agent-based modeling approach for the spread of cutaneous leishmaniasis disease in Central Iran, Isfahan. *Environmental Modelling and Software*, 82, 330–346.
- Rand, W., & Rust, R. T. (2011). Agent-based modeling in marketing: Guidelines for rigor. *International Journal of Research in Marketing*, 28(3), 181–193.
- Reiker, T., Golumbeanu, M., Shattock, A., Burgert, L., Smith, T. A., Filippi, S., Cameron, E., & Penny, M. A. (2021). Emulator-based bayesian optimization for efficient multi-objective calibration of an individual-based model of malaria. *Nature. Communications*, 12(1), 1–11.
- Romero-Mujalli, D., Cappelletto, J., Herrera, E., & Tárano, Z. (2017). The effect of social learning in a small population facing environmental change: An agent-based simulation. *Journal of Ethology*, 35(1), 61–73.
- Rosés, R., Kadar, C., & Malleon, N. (2021). A data-driven agent-based simulation to predict crime patterns in an urban environment. *Computers, Environment and Urban Systems*, 89, 101660.
- Rummery, G. A., & Niranjan, M. (1994). On-line Q-learning using connectionist systems. Technical Report 166, Cambridge University Engineering Department.
- Russell, S., & Norvig, P. (2020). Artificial intelligence: A modern approach. Pearson series in artificial intelligence (pp. 669–675). Pearson.
- Sadeghipour, M., Khoshnevisan, M., Jafari, A., & Shariatpanahi, S. (2017). Friendship network and dental brushing behavior among middle school students: An agent based modeling approach. *PLoS One*, 12(1), e0169236.

- Saeedi, S. (2018). Integrating macro and micro scale approaches in the agent-based modeling of residential dynamics. *International Journal of Applied Earth Observation and Geoinformation*, 68, 214–229.
- Salas-Rodríguez, D., & Rivas-Tovar, L. A. (2022). Agent-based modelling for evaluation of transportation mode selection in the state of Guanajuato, Mexico. *Computación Sistemas*, 26(4), 1–13.
- Salle, I. (2015). Modeling expectations in agent-based models—An application to central bank's communication and monetary policy. *Economic Modelling*, 46, 130–141.
- Sánchez-Marño, N., Alonso-Betanzos, A., Fontenla-Romero, O., Brinquis-Núñez, C., Polhill, J., Craig, T., Dumitru, A., & García-Mira, R. (2015). An agent-based model for simulating environmental behavior in an educational organization. *Neural Processing Letters*, 42(1), 89–118.
- Schauder, S., Thomsen, M., & Nayga, R., Jr. (2020). Agent-based modeling insights into the optimal distribution of the fresh fruit and vegetable program. *Preventive Medicine Reports*, 20, 101173.
- Schwarz, C. V., Reiser, B. J., Davis, E. A., Kenyon, L., Achér, A., Fortus, D., Schwartz, Y., Hug, B., & Krajcik, J. (2009). Developing a learning progression for scientific modeling: Making scientific modeling accessible and meaningful for learners. *Journal of Research in Science Teaching*, 46(6), 632–654.
- Scott, R., MacPherson, B., & Gras, R. (2018). A comparison of stable and fluctuating resources with respect to evolutionary adaptation and life-history traits using individual-based modeling and machine learning. *Journal of Theoretical Biology*, 459, 52–66.
- Sculley, D. (2010). Web-scale k-means clustering. In *Proceedings of the 19th international conference on world wide web—WWW*. ACM Press.
- Sert, E., Bar-Yam, Y., & Morales, A. (2020). Segregation dynamics with reinforcement learning and agent based modeling. *Scientific Reports*, 10(1), 11771.
- Shiono, T. (2021). Estimation of agent-based models using bayesian deep learning approach of bayesflow. *Journal of Economic Dynamics and Control*, 125, 104082.
- Siegfried, R. (2014). *Modeling and simulation of complex systems*. Springer Fachmedien Wiesbaden.
- Snyder, M. N., Schumaker, N. H., Ebersole, J. L., Dunham, J. B., Comeleo, R. L., Keefer, M. L., Leinenbach, P., Brookes, A., Cope, B., Wu, J., Palmer, J., & Keenan, D. (2019). Individual based modeling of fish migration in a 2-d river system: Model description and case study. *Landscape Ecology*, 34(4), 737–754.
- Solinas, F. M., Bottaccioli, L., Guelpa, E., Verda, V., & Patti, E. (2021). Peak shaving in district heating exploiting reinforcement learning and agent-based modelling. *Engineering Applications of Artificial Intelligence*, 102, 104235.
- Sutton, R. S., & Barto, A. G. (2018). *Reinforcement learning*. The MIT Press.
- Taghikhah, F., Voinov, A., Filatova, T., & Polhill, J. G. (2022). Machine-assisted agent-based modeling: Opening the black box. *Journal of Computational Science*, 64, 101854.
- Tarvid, A. (2016). Complex adaptive systems and agent-based modelling. *Agent-Based Modelling of Social Networks in Labour-Education Market System*, 3, 23–38.
- ten Broeke, G., van Voorn, G., Ligtenberg, A., & Molenaar, J. (2021). The use of surrogate models to analyse agent-based models. *Journal of Artificial Societies and Social Simulation*, 24(2), 3.
- Tesfatsion, L. (2002). Agent-based computational economics: Growing economies from the bottom up. *Artificial Life*, 8(1), 55–82.
- van der Hoog, S. (2019). Surrogate modelling in (and of) agent-based models: A prospectus. *Computational Economics*, 53(3), 1245–1263.
- van der Vaart, E., Beaumont, M. A., Johnston, A. S., & Sibly, R. M. (2015). Calibration and evaluation of individual-based models using approximate Bayesian computation. *Ecological Modelling*, 312, 182–190.
- van der Vaart, E., Johnston, A. S., & Sibly, R. M. (2016). Predicting how many animals will be where: How to build, calibrate and evaluate individual-based models. *Ecological Modelling*, 326, 113–123.
- van Strien, M. J., Huber, S. H., Anderies, J. M., & Grêt-Regamey, A. (2019). Resilience in social-ecological systems: Identifying stable and unstable equilibria with agent-based models. *Ecology and Society*, 24(2), 1–24.
- Wang, C., Hu, M., Yang, L., & Zhao, Z. (2021). Prediction of air traffic delays: An agent-based model introducing refined parameter estimation methods. *PLoS ONE*, 16, 1–22.
- Watkins, C. J. C. H. (1989). *Learning from delayed rewards*. PhD thesis. King's College.
- Watson, J. W., Boyd, R., Dutta, R., Vasdekis, G., Walker, N. D., Roy, S., Everitt, R., Hyder, K., & Sibly, R. M. (2022). Incorporating environmental variability in a spatially-explicit individual-based model of european sea bass. *Ecological Modelling*, 466, 109878.
- White, L., Basurra, S., Gaber, M. M., Alsewari, A. A., Saeed, F., & Addanki, S. M. (2022). Agent-based simulations using genetic algorithm calibration: A children's services application. *IEEE Access*, 10, 88386–88397.
- Wilensky, U., & Rand, W. (2015). *An introduction to agent-based modeling: Modeling natural, social, and engineered complex systems with NetLogo*. MIT Press.
- Witten, I. H., Frank, E., Hall, M. A., & Pal, C. J. (2017). *Data mining: Practical machine learning tools and techniques* (4th ed.). Morgan Kaufmann.
- Wolf, I., Schröder, T., Neumann, J., & de Haan, G. (2015). Changing minds about electric cars: An empirically grounded agent-based modeling approach. *Technological Forecasting and Social Change*, 94, 269–285.
- Xiao, S., & Liu, R. (2021). Studies of covid-19 outbreak control using agent-based modeling. *Complex Systems*, 30(3), 297–321.
- Xie, S., Lawniczak, A. T., & Gan, C. (2022). Optimal number of clusters in explainable data analysis of agent-based simulation experiments. *Journal of Computational Science*, 62, 101685.
- Xu, S., Chen, X., Xie, J., Rahman, S., Wang, J., Hui, H., & Chen, T. (2019). Agent-based modeling and simulation of the electricity market with residential demand response. *CSEE Journal of Power and Energy Systems*, 7, 99.
- Xu, T., Gao, J., Coco, G., & Wang, S. (2020). Urban expansion in Auckland, New Zealand: A gis simulation via an intelligent self-adapting multiscale agent-based model. *International Journal of Geographical Information Science*, 34(11), 2136–2159.
- Yao, F., Zhu, J., Yu, J., Chen, C., & Chen, X. (2020). Hybrid operations of human driving vehicles and automated vehicles with data-driven agent-based simulation. *Transportation Research Part D: Transport and Environment*, 86, 102469.
- Ye, P., Chen, Y., Zhu, F., Lv, Y., Lu, W., & Wang, F. (2021). Bridging the micro and macro: Calibration of agent-based model using mean-field dynamics. *IEEE Transactions on Cybernetics*, 52, 11397–11406.
- Zhang, H., Vorobeychik, Y., Letchford, J., & Lakkaraju, K. (2016). Data-driven agent-based modeling, with application to rooftop solar adoption. *Autonomous Agents and Multi-Agent Systems*, 30(6), 1023–1049.
- Zhang, W., Valencia, A., & Chang, N.-B. (2021). Synergistic integration between machine learning and agent-based modeling: A multidisciplinary review. *IEEE Transactions on Neural Networks and Learning Systems*, 34, 1–21.



- Zhang, Y., Grignard, A., Lyons, K., Aubuchon, A., & Larson, K. (2018). *Real-time machine learning prediction of an agent-based model for urban decision-making (extended abstract)* (Vol. 3, pp. 2171–2173). International Foundation for Autonomous Agents and Multiagent Systems (IFAAMAS).
- Zhang, Y., Li, Z., & Zhang, Y. (2020). Validation and calibration of an agent-based model: A surrogate approach. *Discrete Dynamics in Nature and Society*, 9, 1–10.
- Zornić, N., & Marković, A. (2022). A methodological framework for the integration of machine learning algorithms into agent-based simulation models. *JUCS—Journal of Universal Computer Science*, 28(5), 540–562.

AUTHOR BIOGRAPHIES

Alejandro Platas-López is currently a doctoral candidate in Artificial Intelligence at the Universidad Veracruzana and a lecturer at Universidad Anáhuac Online, focusing on AI. His research interests include evolutionary computation, machine learning, and agent-based modeling and simulation. His Ph.D. work follows an MSc in AI and an undergraduate degree in Economics, both from Universidad Veracruzana. He has received recognition from the COVEICYDET and won awards at national and international research competitions.

Alejandro Guerra-Hernández is currently Professor in the Instituto de Investigaciones en Inteligencia Artificial of the Universidad Veracruzana, México. He is member of the Mexican National Research System. He has done visiting stays at the Universitat de València, Spain, and the Université de Paris 6, France. He received a Ph.D., full marks and honors, from the Université de Paris 13, where he was Assistant Professor at the Laboratoire d'Informatique de Paris Nord. He has been PC member of various international conferences including IJCAI, ICMAS, AAMAS, EUMAS, MASTA, EPIA, Iberamia, and MICAI. His main research interests are Learning in Multi-Agent Systems and Agent-Based Modelling/Simulation.

Marcela Quiroz is a Full-Time Researcher with the Artificial Intelligence Research Institute at the Universidad Veracruzana in Xalapa City, Mexico. Her research interests include: combinatorial optimization, metaheuristics, experimental algorithms, characterization and data mining. She received her Ph.D. in Computer Science from the Instituto Tecnológico de Tijuana, Mexico. She studied engineering in computer systems and received the degree of master in computer science at the Instituto Tecnológico de Ciudad Madero, Mexico. She is a member of the Mexican National Researchers System (SNI), and also a member of the directive committees of the Mexican Computing Academy (AMexComp) and the Mexican Robotics Federation (FMR).

Nicandro Cruz-Ramírez is a Full-Time Professor at the Artificial Intelligence Research Institute at Universidad Veracruzana, Mexico. He obtained his Ph.D. in Artificial Intelligence from the University of Sheffield, an MSc in AI from Universidad Veracruzana, and a Bachelor's degree in Electronic Engineering from the Instituto Tecnológico de Orizaba. Throughout his career, he has served in various evaluative roles on academic and research committees. His research primarily focuses on Bayesian networks and data mining.

How to cite this article: Platas-López, A., Guerra-Hernández, A., Quiroz-Castellanos, M., & Cruz-Ramírez, N. (2023). A survey on agent-based modelling assisted by machine learning. *Expert Systems*, e13325. <https://doi.org/10.1111/exsy.13325>

CAPÍTULO 4

APRENDIZAJE OFFLINE EN MODELACIÓN BASADA EN AGENTES

Contenido del capítulo

4.1. Propuesta de documentación	113
4.1.1. Aprendizaje fuera de línea	114
4.1.2. A priori	117
4.1.3. A posteriori	117
4.2. Caso de estudio: Evasión fiscal	118
4.2.1. Aprendizaje supervisado a priori	122
4.2.2. Aprendizaje supervisado a posteriori	124
4.3. Apéndice	132

La necesidad de integrar técnicas más sofisticadas para mejorar el diseño, validación y análisis de modelos ha emergido como una respuesta directa a las críticas y desafíos inherentes al campo. Aunque las potencialidades del aprendizaje automático en la simulación basada en agentes son vastas, aún no existe un consenso claro o un camino definido para su aplicación, lo que a menudo puede llevar a esfuerzos fragmentados y resultados no óptimos. El presente capítulo se adentra en el mundo del aprendizaje offline en ABM, proporcionando claridad y estructura a este campo emergente.

La primera sección delinearé inicialmente el contexto general de cómo el aprendizaje offline se posiciona dentro de la documentación del protocolo ODD, sin importar el objetivo de su implementación. Posteriormente, se profundiza en los detalles de cómo este aprendizaje

La segunda sección sirve como una ilustración práctica de los conceptos metodológicos propuestos. A través de este caso, se examina la forma en que el aprendizaje supervisado puede ser encapsulado dentro de una simulación, asistiendo en la tarea de diseño del modelo. Posteriormente, mostramos cómo el aprendizaje offline también permite analizar y entender los resultados de una manera más profunda y matizada. Se abordan los detalles técnicos y metodológicos de este proceso, ofreciendo una guía paso a paso para investigadores y profesionales interesados.

Como se revisó en el Capítulo 2, el protocolo de Visión General, Conceptos de Diseño y Detalles (ODD, por sus siglas en inglés) (Grimm et al., 2010, 2020), es un documento estandarizado que proporciona una descripción consistente, lógica y legible de la estructura y dinámicas de los ABM. Consiste en siete elementos, como se muestra en la Figura 4.1. Estos elementos están agrupados en tres niveles de descripción, que varían de general a específico.

		Basic principles
O	1. Purpose and patterns	Emergence
	2. Entities, state variables and scales	Adaptation
	3. Process overview and scheduling	Objetives
D	4. Design concepts	Learning
D	5. Initialization	Prediction
	6. Input data	Sensing
		Interaction
	7. Submodels	Stochasticity
		Collectives
		Observation

Dado que la inclusión de ML en ABM se ha producido gradualmente, el protocolo ODD original [Grimm et al. \(2006\)](#) no consideró el aprendizaje entre los elementos de diseño. En

la primera actualización del protocolo [Grimm et al. \(2010\)](#), se reconoció la relevancia del aprendizaje, incluyéndolo como concepto de diseño, pero adoptando únicamente la modalidad en línea. En lugar de mejorar la especificación del aprendizaje como elemento de diseño, la segunda actualización del protocolo ODD [Grimm et al. \(2020\)](#) se centró en otros aspectos como la claridad, la replicación y el realismo estructural. De hecho, en su estado actual, las preguntas del protocolo ODD sobre el aprendizaje incluyen solo: “¿Muchos individuos o agentes (pero también organizaciones e instituciones) cambian sus rasgos adaptativos a lo largo del tiempo como consecuencia de su experiencia? ¿Si es así, cómo?”.

4.1.1. Aprendizaje fuera de línea

Por lo tanto, la sección de aprendizaje en el protocolo ODD actual pretende describir únicamente las tareas de aprendizaje en línea, dejando de lado la posibilidad de documentar las aplicaciones de aprendizaje fuera de línea, tanto antes como después de la simulación, como se mencionó anteriormente. Parece necesaria una extensión del protocolo ODD que permita la descripción de todas estas tareas de aprendizaje.

En este contexto, esta tesis busca extender las preguntas asociadas con el Protocolo ODD para abordar aspectos adicionales del aprendizaje en la modelación basada en agentes, las cuales se describen a continuación:

¿Cuál es el propósito del aprendizaje?

El aprendizaje puede tener varios propósitos en el contexto de la modelación basada en agentes:

- *Mejorar el desempeño de los agentes:* Siguiendo el enfoque actual del Protocolo ODD, el aprendizaje puede ser utilizado para mejorar el desempeño de los agentes durante la simulación en base a su experiencia previa.
- *Generar datos para simulaciones:* El aprendizaje puede ser utilizado para generar datos sintéticos que serán utilizados en las simulaciones.
- *Inducción de reglas de comportamiento:* Las reglas que definen el comportamiento de los agentes pueden ser inducidas a partir de datos existentes.
- *Calibración de parámetros:* Los parámetros del modelo pueden ser calibrados mediante aprendizaje.
- *Análisis de resultados:* Los resultados de las simulaciones pueden ser analizados utilizando métodos de aprendizaje.

¿Cuándo se realiza el aprendizaje?

El momento en que se realiza el aprendizaje depende de su finalidad:

- *Antes de la simulación:* Para generar datos sintéticos, inducir reglas de comportamiento de agentes o instanciar parámetros del modelo.
- *Durante la simulación:* Para mejorar el desempeño de los agentes en tiempo real, como se espera en el Protocolo ODD actual.
- *Después de la simulación:* Para analizar los datos producidos por la simulación o calibrar parámetros del modelo.

¿Qué componentes del ABM se ven afectados por el aprendizaje?

Varios componentes del ABM pueden ser afectados por el aprendizaje:

- *Agentes*: En el Protocolo ODD actual, los agentes son los únicos componentes que realizan y se ven afectados por el aprendizaje. El aprendizaje puede ayudar a los agentes a mejorar su desempeño durante la simulación.
- *Entorno*: El entorno puede explotar el uso de métodos de aprendizaje automático, por ejemplo, consumiendo datos sintéticos generados por modelos aprendidos.
- *Parámetros del modelo*: Es posible aprender los valores de los parámetros del ABM para obtener un comportamiento deseado de todo el sistema.

¿Cómo se calcula el aprendizaje?

El tipo de algoritmo utilizado y sus datos de entrada dependen del momento y propósito del aprendizaje.

Momento del aprendizaje.

- *Aprendizaje en línea*: Requiere algoritmos de aprendizaje incremental, como los incluidos en el aprendizaje por refuerzo (RL).
- *Aprendizaje fuera de línea*: Las técnicas supervisadas se pueden utilizar tanto antes como después de la simulación.

Propósito del aprendizaje.

- *Calibración de parámetros*: Los enfoques evolutivos son muy adecuados para calibrar parámetros.
- *Explicabilidad*: La preferencia por métodos explicables versus técnicas de caja negra debe estar relacionada con el propósito del aprendizaje.

Si bien las preguntas planteadas anteriormente proporcionan una visión general del aprendizaje fuera de línea en la modelación basada en agentes, es necesario profundizar en ciertas secciones del protocolo ODD para ofrecer una descripción más detallada y específica del proceso de aprendizaje. Esta descripción detallada permitirá identificar con precisión en qué etapa del ciclo de desarrollo se está incidiendo. En las siguientes subsecciones, se analizan consideraciones adicionales que ayudan a definir el momento en el que se realiza el aprendizaje y los elementos a tener en cuenta dentro del protocolo ODD.

4.1.2. A priori

El ML a priori implica el uso de datos existentes relacionados con el fenómeno a simular. El uso de dichos datos debe incluirse en la sección de detalles del Protocolo ODD, concretamente en el elemento denominado “Datos de entrada”, en el que se responde a la siguiente pregunta:

6 Input data

Actual: ¿El modelo usa entradas de fuentes externas como archivos de datos u otros modelos para representar procesos que cambian con el tiempo?

Propuesta: ¿El modelo usa datos de fuentes externas (archivos de datos u otros modelos) para representar **algún elemento en el modelo?** ¿Se modelan los datos con aprendizaje automático?

4.1.3. A posteriori

A posteriori ML implica el uso de datos resultantes de la simulación. El uso de dichos datos debe estar incluido en el apartado de conceptos de diseño del Protocolo ODD, concretamente en el elemento denominado “Observación”, en el que se responde a la siguiente pregunta:

5.11 Observation

Actual: ¿Qué datos se recopilan del ABM para probarlos, comprenderlos y analizarlos, y cómo y cuándo se recopilan? ¿Se usan libremente todos los datos de salida, o solo se muestrean y usan ciertos datos para imitar lo que se puede observar en un estudio empírico?

Propuesta: ¿Qué datos se recopilan del ABM para **calibrar**, probar, comprender y analizar, y cómo y cuándo se recopilan? ¿Se usan libremente todos los datos de salida, o solo se muestrean y usan ciertos datos para imitar lo que se puede observar en un estudio empírico? ¿Se modelan los datos recopilados con aprendizaje automático?

Las preguntas reescritas de esta manera vinculan la fuente de los datos con la sección de aprendizaje en la que se describe la tarea de aprendizaje a realizar. Finalmente, se debe

proporcionar una descripción completa del modelo ML en los submodelos del protocolo ODD, incluida la preparación, el modelado y la evaluación de datos.

4.2. Caso de estudio: Evasión fiscal

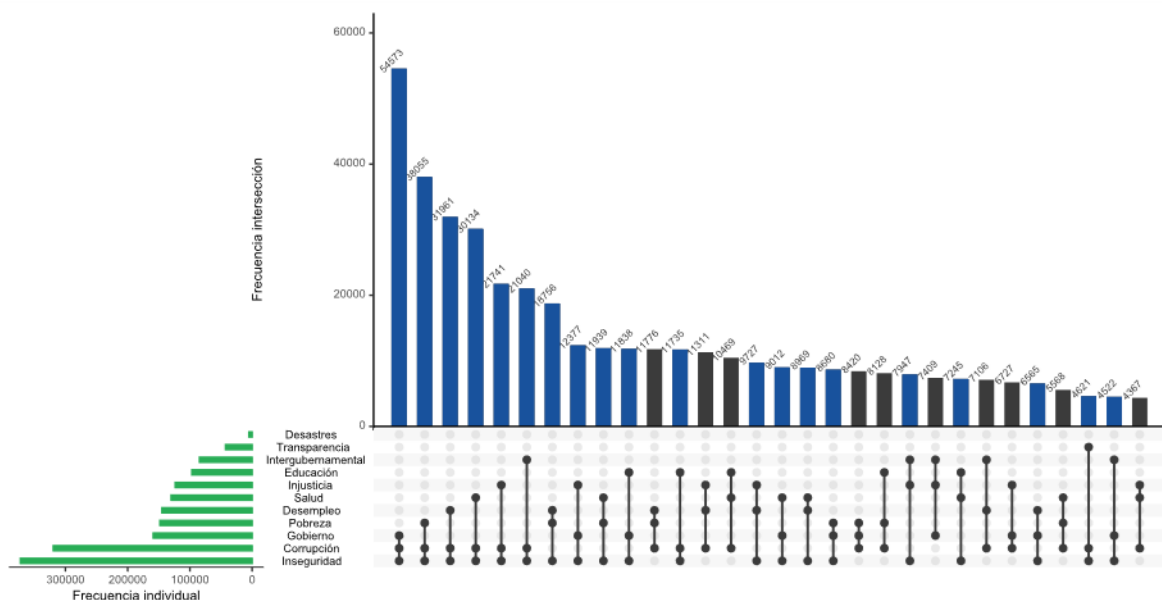
La evasión fiscal, una actividad intencionalmente ilegal, compromete el adecuado funcionamiento de un estado, al reducir los recursos públicos destinados a la provisión de servicios públicos esenciales como la educación, la salud y la seguridad (Alm, 2011; Slemrod y Yitzhaki, 2000). La labor de explorar y entender esta problemática comenzó a formalizarse en la década de los setenta con los trabajos de Allingham y Sandmo (1972) quienes, a partir de la economía del delito (Becker, 1968), establecieron un acercamiento teórico para examinar el comportamiento de declaración de impuestos.

Respecto a los factores económicos, Skinner y Slemrod (1985) señalan que el diseño y la implementación de un sistema tributario complejo y con altas tasas impositivas pueden motivar la evasión con el propósito de reducir las cargas impositivas. Por otro lado, la percepción de la aplicación de penalizaciones y la detección de la evasión por parte de las autoridades pueden verse afectadas por el contexto sociocultural, por ejemplo, la aceptación social de la evasión y la percepción de corrupción (Sandmo, 2005) o bien, la percepción de un sistema tributario inequitativo.

En este sentido, las penalizaciones y la detección son dos herramientas clave en la lucha contra la evasión fiscal. En tanto la efectividad de las penalizaciones como política de sustitución depende de factores como la aceptación social y la percepción de cumplimiento por parte de otros contribuyentes (Sandmo, 2005). En el contexto de México, la percepción de inseguridad y corrupción destacan como problemas críticos que los ciudadanos identifican frecuentemente, tal como se evidencia en la Figura 4.2. Estas percepciones pueden atenuar el efecto disuasorio de las penalizaciones y reducir la eficacia de los esfuerzos de cumplimiento, lo que hace que este análisis multidimensional adquiera una relevancia particular en el país.

En cuanto a los factores culturales e institucionales, diversos autores (Alm y Martínez-Vazquez, 2007; Basilgan y Christiansen, 2016; Aguirre-Quezada y Sánchez-Ramírez, 2019) advierten la influencia de éstos en la decisión de las personas contribuyentes para evadir impuestos. Por un lado, el enfoque de Alm y Martínez-Vazquez enfatiza en la moral tributaria y la confianza en las instituciones, y por el otro, Aguirre-Quezada y Sánchez-Ramírez sobre

Figura. 4.2: Prevalencia de Inseguridad y Corrupción como Principales Problemas Percibidos en los Estados.



Fuente: Elaboración propia.

las consecuencias sociales y la inequidad, sugieren que la evasión fiscal está incrustada en un entramado de relaciones y percepciones sociales.

Desde un enfoque holístico, [Khlif y Achek \(2015\)](#) presentan una amplia gama de determinantes de la evasión fiscal, que incluyen aspectos demográficos, culturales, de comportamiento, legales, institucionales y económicos. Esto hace eco de los hallazgos previos de [Daude et al. \(2012\)](#), que también reconocen la multiplicidad de factores que convergen en la evasión fiscal, particularmente en países en desarrollo como México. Identificaron dos grupos de factores: factores socioeconómicos, como la edad o la educación, y factores institucionales, como la confianza en el gobierno y la satisfacción con la calidad de los servicios públicos.

En su trabajo de investigación, [Aguirre-Quezada y Sánchez-Ramírez \(2019\)](#) hacen énfasis en el marco jurídico que regula la evasión fiscal en México y las consecuencias sociales que esto genera, como la desigualdad y la inequidad. Señalan que la evasión fiscal resulta en una carga injusta para aquellos que sí cumplen con sus obligaciones tributarias, lo que contribuye a una injusticia en la sociedad y afecta las finanzas públicas debido a una menor recaudación.

Aquí podemos ver un punto de convergencia entre [Aguirre-Quezada y Sánchez-Ramírez](#)

(2019), y una investigación previa (Platas-López, 2014). Mientras que los primeros autores señalan cómo la evasión fiscal contribuye a la injusticia social y afecta las finanzas públicas, en Platas-López (2014) se evidenció que la eficiencia en la recaudación del ISN es desigual entre diferentes regiones, lo que puede estar exacerbando los problemas de inequidad señalados por Aguirre-Quezada y Sánchez-Ramírez. Esto sugiere que mejorar la eficiencia en la recaudación del ISN podría ser una forma de abordar la inequidad en la distribución de la carga fiscal.

Por otro lado, Ayala-Gaytán et al. (2016) abordan el impacto económico del ISN en los salarios y el empleo. Encuentran que el ISN tiene un impacto negativo en los salarios, pero no en el empleo, y que la carga del impuesto recae principalmente en los trabajadores. Este último hallazgo es especialmente relevante cuando se lo considera junto con el trabajo de Aguirre-Quezada y Sánchez-Ramírez (2019). Si la carga del ISN recae desproporcionadamente sobre los trabajadores, esto podría ser una fuente de la iniquidad como consecuencia de la evasión fiscal.

En conjunto, estos estudios sugieren un panorama complejo en el que la evasión fiscal y la ineficiencia en la recaudación de impuestos están interrelacionadas y tienen implicaciones significativas en términos de justicia social y equidad. Esto subraya la importancia de abordar estos problemas de manera integral, considerando tanto el marco legal como la eficiencia en la recaudación de impuestos y las consecuencias económicas y sociales.

Aunque las primeras investigaciones teóricas proporcionan un punto de partida para evaluar suposiciones sobre la evasión fiscal, el método básico de las fórmulas clásicas, es decir, el uso de funciones de utilidad y las suposiciones de homogeneidad y racionalidad del contribuyente, describen de manera insuficiente los comportamientos de los contribuyentes observados en la práctica.

Considerando estas perspectivas, el uso de modelos basados en agentes, como sugieren Mittone y Patelli (2000), Davis et al. (2003), y Hokamp (2014), se convierte en una metodología especialmente valiosa. Ya que estas investigaciones adoptan enfoques particulares para examinar este fenómeno desde diversos ángulos y aportan una serie de perspectivas interconectadas, aunque distintas, sobre cómo la psicología del contribuyente, las normas sociales, la aplicación de la ley y la heterogeneidad de las poblaciones afectan la evasión fiscal.

Respecto a la psicología del contribuyente y los beneficios de los bienes públicos en su modelo de evasión fiscal, [Mittone y Patelli \(2000\)](#) identifican diferentes clases de contribuyentes, cada una con una función de utilidad única. Aquí se plantea un punto de inflexión en la evasión fiscal, donde se introduce el papel de las auditorías en la creación de bienes públicos, que, a su vez, informan la utilidad percibida por los contribuyentes.

Por su parte, [Davis et al. \(2003\)](#) desplazan el foco hacia las normas sociales y la aplicación de la ley. Su modelo de evasión fiscal pone de manifiesto la interacción dinámica entre estas dos fuerzas. Concluyen que el impacto de la aplicación de la ley en el cumplimiento de los impuestos es dependiente del estado inicial de la población contribuyente, revelando que las diferentes poblaciones pueden responder de manera única a las políticas fiscales similares. Su trabajo ofrece una perspectiva para entender la variación geográfica y temporal en la evasión de impuestos, un aspecto relevante en el contexto de la inseguridad y la corrupción en México.

Mientras tanto [Hokamp \(2014\)](#), aporta una nueva dimensión a esta discusión al considerar la heterogeneidad de las poblaciones y las interacciones sociales. Su inclusión del envejecimiento en el modelo introduce la idea de que las normas sociales se actualizan a lo largo del ciclo de vida de una persona contribuyente. Al explorar la provisión de bienes públicos, el estudio respalda la idea de que dicha heterogeneidad puede conducir a fluctuaciones en la evasión de impuestos.

Estos estudios ofrecen una serie de perspectivas que nos ayudaran a formar el presente caso de estudio. La interacción de la psicología del contribuyente, las normas sociales y la aplicación de la ley, junto con las influencias de la heterogeneidad de la población y las fluctuaciones en el tiempo y el espacio, pueden proporcionar un marco sólido para explorar el papel de los determinantes institucionales, como la inseguridad y la corrupción, en la evasión del impuesto a la nómina en México. En particular, los diferentes enfoques de modelización de estos estudios permitirán comprender cómo estos factores pueden influir en el comportamiento de las personas contribuyentes.

Tras haber analizado la relevancia y la contribución de diversos estudios en el ámbito de la evasión fiscal, se hace necesario abordar la modelación basada en agentes como una herramienta potente para analizar y comprender la evasión del ISN en México, y cómo diferentes determinantes micro y macro pueden incidir en la evasión.

Una consideración importante en el desarrollo de ABM es cómo aprovechar la existencia y generación de datos para asistir en el diseño y análisis del modelo respectivamente. El ML ha emergido como una herramienta valiosa en este aspecto. No solo permite incorporar datos reales en la fase de diseño para crear modelos más precisos, sino que también se puede utilizar en la fase de análisis para identificar patrones y relaciones complejas que podrían no ser evidentes a través de análisis convencionales.

Dicho esto, el análisis de un modelo basado en agentes puede ser complicado debido a la naturaleza altamente dinámica y no lineal que suelen presentar estos modelos. El gran número de interacciones y la emergencia de comportamientos inesperados requieren herramientas y enfoques analíticos sofisticados. Para abordar estos desafíos, en las siguientes secciones se presentan dos enfoques distintos que utilizan aprendizaje supervisado para asistir en diferentes etapas del proceso de modelado basado en agentes.

4.2.1. Aprendizaje supervisado a priori

En esta subsección, se explora cómo el aprendizaje supervisado puede ser utilizado en la fase de diseño de un modelo basado en agentes. Utilizando datos reales, se instanciará un modelo que los agentes usarán en la simulación para tomar decisiones sobre si evadir o no impuestos. Los detalles de la implementación descritos siguiendo el protocolo ODD se presentan en el Anexo A. La figura 4.3 muestra de manera resumida el proceso de integración de ML.

Con la propuesta de extensión para la inclusión del aprendizaje se responde a las siguientes interrogantes.

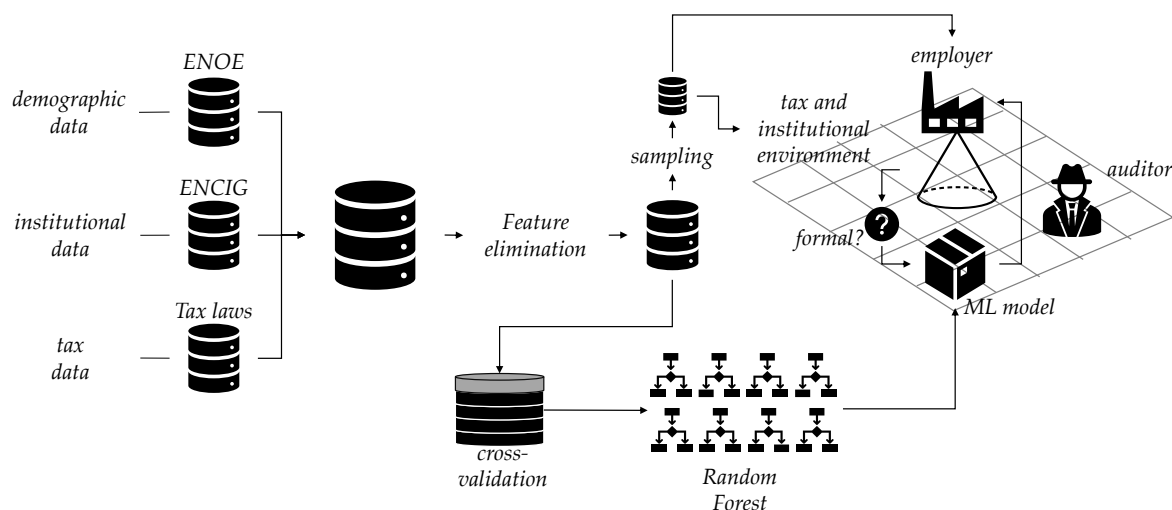
¿Cuál es el propósito del aprendizaje?

El propósito del aprendizaje es inducir las reglas que definen el comportamiento de los agentes a partir de los datos existentes.

¿Cuándo se realiza el aprendizaje?

El aprendizaje se realiza antes de la simulación para obtener modelos para generar las reglas de los agentes incluidos en el ABM.

Figura. 4.3: Diagrama de integración de ML en el Modelo de Evasión de Impuestos sobre Nómina.



Fuente: Elaboración propia.

¿Qué componentes del ABM se ven afectados por el aprendizaje?

El entorno se ve afectado al encapsular en un modelo aprendido los datos del sistema fiscal y la calidad de los bienes públicos.

¿Cómo se calcula el aprendizaje?

El aprendizaje se calcula mediante un algoritmo supervisado.

6 Input data

¿El modelo usa datos de fuentes externas (archivos de datos u otros modelos) para representar **algún elemento en el modelo**?

Sí, el modelo utiliza datos externos del sistema tributario mexicano para instanciar el ambiente en el que interactúan los agentes.

¿Se modelan los datos con aprendizaje automático? Si, se modelaron con una implementación de Random Forest del paquete R "ranger"(Wright y Ziegler, 2017) para aprender el ambiente a partir de los datos disponibles.

4.2.2. Aprendizaje supervisado a posteriori

En esta parte, se examina cómo el aprendizaje supervisado puede ser útil para analizar el modelo después de su implementación. Se empleará una red bayesiana para construir un modelo transparente, cuyos parámetros estarán optimizados con algoritmos de evolución diferencial. Este enfoque busca mejorar una posible desventaja de estos modelos: la precisión en la clasificación, la cual pueda ser una de las causas por las cuales la comunidad suele emplear algoritmos de caja negra.

Siguiendo nuestra propuesta para extender y enriquecer la documentación bajo el protocolo ODD, se considera esencial responder a ciertas preguntas clave que faciliten una comprensión más profunda del modelo en su conjunto. Esta interrogante busca abordar la integración del aprendizaje automático en el esquema de simulación y su relevancia en el análisis del sistema. A continuación, presentamos y respondemos a esta pregunta:

5.11 Observation

¿Se modelan los datos recopilados con aprendizaje automático?

Si, se modelan con una red bayesiana para aprender cómo distintas variables del sistema influyen en la Extensión de la Evasión Fiscal.

Dentro de los enfoques para aprender la estructura de una red bayesiana, existen diferentes heurísticas eficientes [Scanagatta et al. \(2019\)](#). Algunas de ellas guían la búsqueda aliviando la suposición de independencia naïve entre los atributos, como es el caso del algoritmo Tree Augmented Naïve Bayes (TAN) [Friedman et al. \(1997\)](#). TAN es un método semi-naïve de aprendizaje bayesiano que utiliza una estructura de árbol en la que cada atributo solo depende de la clase y otro atributo.

Para realizar la clasificación, se utiliza un árbol de expansión ponderado máximo que maximiza la distribución de probabilidad subyacente de los datos de entrenamiento [Chow y Liu \(1968\)](#). Este algoritmo emplea una métrica (como la longitud de descripción mínima, MDL) para puntuar la red bayesiana durante la búsqueda. La búsqueda comienza con la red vacía y aplica sucesivamente operaciones locales que mejoran al máximo la puntuación hasta que se encuentra un mínimo local. Las operaciones aplicadas por el procedimiento de búsqueda incluyen la adición de arcos, la eliminación de arcos y la inversión de arcos.

Como clasificador, TAN mejora al Naïve Bayes en términos de sesgo a costa de un aumento en la varianza debido al crecimiento en el número de parámetros a estimar. Aunque la red aprendida con TAN generalmente supera al Naïve Bayes, el primero tiene un peor desempeño al lidiar con conjuntos de datos con muchos atributos. [Friedman et al. \(1997\)](#) explica este comportamiento dada la definición de MDL y otras métricas similares, es decir, midiendo el error de la red aprendida sobre todas las variables en el dominio, incluida la variable de clase. Sin embargo, minimizar este error no necesariamente minimiza el error local al predecir la variable de clase dada los atributos, especialmente en casos donde hay muchos atributos.

Aunque no se define explícitamente qué se considera un conjunto de datos con “muchos atributos”, la naturaleza de la simulación basada en agentes llevada a cabo en nuestro caso de estudio nos llevó a considerar que los resultados de la salida de la simulación poseen la complejidad y dimensionalidad adecuadas para beneficiarse del uso del algoritmo TAN.

Los resultados de nuestra simulación se componen de parámetros que reflejan cambios en el contexto de la evasión fiscal en México. Estos parámetros han sido derivados considerando determinantes institucionales como la calidad en la provisión de bienes públicos. En este estudio, se ha dado particular énfasis a dos determinantes institucionales clave según la encuesta nacional de calidad e impacto gubernamental: la percepción de inseguridad y la percepción de corrupción.

El conjunto de datos resultante de la simulación, refleja los siguientes parámetros:

- π Tasa de penalización: Esta variable refleja la sanción que un individuo podría enfrentar al evadir impuestos.
- α Probabilidad de auditoría: Indica la probabilidad de que un individuo o entidad sea sometido a una revisión de sus finanzas por parte de las autoridades fiscales.
- $\Delta\theta$ Variación en la tasa impositiva: Representa las fluctuaciones en la tasa impositiva a la que están sujetos los contribuyentes.
- ϵ_{TC} Eficiencia de la recolección de impuestos: Refleja la efectividad con la que las autoridades fiscales recaudan los impuestos.

ϵ_{AP} Eficiencia en el proceso de auditoría: Indica la eficacia con la que las autoridades fiscales llevan a cabo las auditorías.

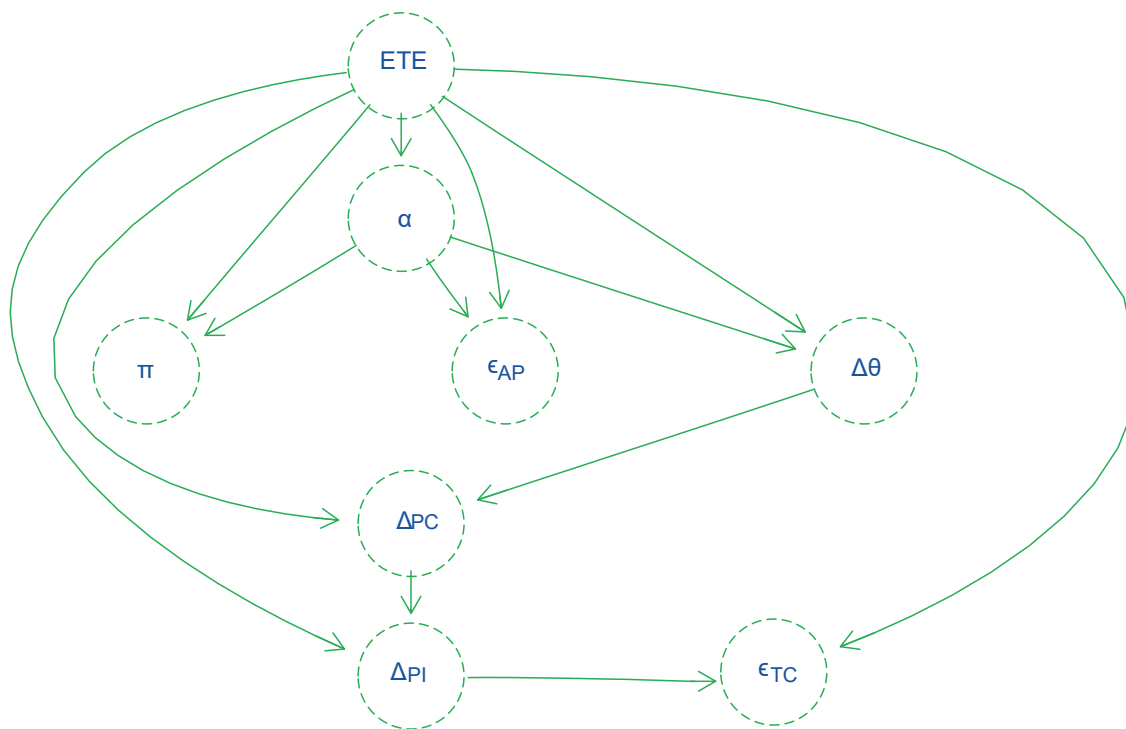
ΔPI Variación en la percepción de inseguridad: Representa cambios en la percepción pública sobre la inseguridad en el estado.

ΔPC Variación en la percepción de corrupción: Representa cambios en la percepción pública sobre la corrupción en el estado.

Por último, la variable respuesta, ETE (Extension of Tax Evasion), sirve como un indicador de la extensión o magnitud de la evasión fiscal en el conjunto de datos.

Dado que nuestra simulación genera datos con siete parámetros influyentes y una variable objetivo, es imperativo adoptar una técnica que pueda capturar adecuadamente las complejas relaciones subyacentes. Es aquí donde el algoritmo TAN se presenta como una opción robusta, capaz de descomponer estas relaciones en una estructura que mitiga la suposición de independencia entre atributos y ofrece una representación más transparente y comprensible del fenómeno estudiado. La estructura resultante se puede apreciar en la Figura 4.4.

Figura. 4.4: Estructura de la red bayesiana aprendida con el algoritmo Tree Augmented Naive Bayes (TAN) (Chow y Liu, 1968).



Fuente: Elaboración propia.

La probabilidad de auditoría, representada por α , parece aumentar en respuesta a una alta Extensión de la Evasión Fiscal (ETE). Esta relación sugiere un mecanismo natural para desincentivar la evasión, ya que es lógico pensar que ante un incremento en la evasión fiscal, las autoridades podrían intensificar las auditorías. Además, a medida que esta probabilidad de auditoría aumenta, también parece aumentar la tasa de penalización (π) para los evasores. Esta correlación puede reflejar una estrategia de endurecimiento de las sanciones frente a un panorama donde la evasión se torna más frecuente. Interesantemente, la eficiencia en el proceso de auditoría, ϵ_{AP} , también parece estar correlacionada con α , lo que podría sugerir que a medida que aumenta la probabilidad de auditoría, se asignan más recursos o se adoptan técnicas más eficientes en el proceso.

Por otro lado, la tasa de penalización, π , no solo está influenciada por α , sino también directamente por ETE. Si hay una alta Extensión de la Evasión Fiscal, es plausible que las autoridades decidan incrementar las penalizaciones como una medida disuasoria. Esta relación parece subrayar la importancia de adaptar las sanciones al entorno fiscal actual.

Un aspecto intrigante es cómo la probabilidad de auditoría y la Extensión de la Evasión Fiscal influyen en $\Delta\theta$, que representa un cambio en la tasa impositiva. Estos vínculos sugieren que, en un entorno con alta evasión o alta probabilidad de auditoría, podría haber reconsideraciones en la política impositiva. Estos cambios en la tasa impositiva, a su vez, parecen tener un impacto en la percepción de corrupción, ΔPC . Esto sugiere que, si el público percibe que las tasas impositivas están cambiando en respuesta a la evasión o de manera arbitraria, esto podría aumentar la percepción de corrupción.

Además, esta variación en la percepción de corrupción influencia ΔPI , una medida de la percepción de inseguridad. Esto podría reflejar la idea de que una sociedad que se percibe como corrupta también puede ser vista como insegura. Finalmente, esta percepción de inseguridad tiene un efecto sobre la eficiencia en la recolección de impuestos, ϵ_{TC} . Esto sugiere que en lugares donde las personas se sienten inseguras, podrían ser menos propensas a cumplir con sus obligaciones fiscales, lo que a su vez afectaría la eficiencia en la recolección.

La red bayesiana propuesta ilustra una serie de relaciones interconectadas que reflejan cómo diferentes aspectos del sistema fiscal y las percepciones públicas pueden interactuar entre sí en el contexto de la evasión fiscal. Es crucial, no obstante, validar estas relaciones con expertos en el tema para una interpretación más precisa y fundamentada.

Si bien el modelo propuesto a partir de la red bayesiana brinda una estructura comprensible y transparente, afronta el reto de optimizar su precisión y capacidad de generalización. Estas deficiencias pueden ser una razón por la cual muchos investigadores se han inclinado hacia modelos de “caja negra”. Para contrarrestar esto y potenciar la exactitud de modelos transparentes, se propone la aplicación del aprendizaje discriminativo de los parámetros, específicamente mediante la evolución diferencial (Price y Storn, 1996) para optimizar el Conditional Log Likelihood (CLL).

De acuerdo con lo reportado en [Platas-López et al. \(2021\)](#), el aprendizaje discriminativo de los parámetros de una red bayesiana es una alternativa prometedora. Al optimizar el CLL mediante la Evolución Diferencial, se busca aumentar la precisión del modelo al minimizar el error de clasificación, favoreciendo así tanto la precisión como la generalización del modelo.

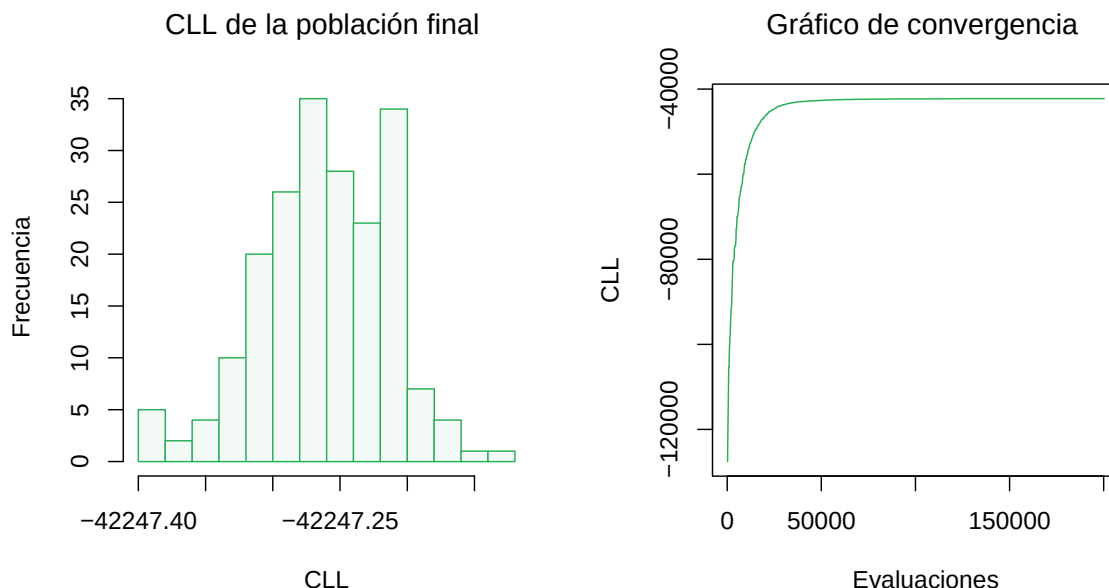
El software R presentado en [Platas-López et al. \(2023a\)](#) permite optimizar los parámetros

a través de variantes de Evolución Diferencial basadas en JADE (Zhang y Sanderson, 2009), como L-SHADE (Tanabe y Fukunaga, 2014), las cuales destacan particularmente por sus mecanismos de adaptación de parámetros. Estos adaptan de manera dinámica los parámetros según las necesidades del problema, evitando la tediosa tarea de calibración manual y, al mismo tiempo, incrementando la capacidad del modelo para adaptarse y generalizar en diferentes escenarios.

Para poder optimizar los parámetros de la estructura de la red aprendida con TAN, se ha optado por emplear la versión de JADE con Archivo. Esta elección no es aleatoria: la variante JADE de la Evolución Diferencial ha demostrado ser especialmente eficaz en escenarios de optimización complejos, especialmente debido a sus mecanismos de adaptación de parámetros. Estos mecanismos permiten que JADE se ajuste dinámicamente a la naturaleza del problema sin requerir una laboriosa calibración manual, lo que resulta ser crucial para garantizar tanto la precisión como la capacidad de generalización del modelo.

El Archivo de JADE conserva un conjunto determinado de soluciones previamente descartadas, reintroduciéndolas de manera aleatoria durante el proceso evolutivo para fomentar la diversidad dentro de la población. Hemos definido el tamaño de la población como $NP = 200$ y un criterio de paro que se activa tras $G = 1,000$ generaciones o cuando el algoritmo no registre mejoras a lo largo de 10 generaciones consecutivas. Esta configuración busca encontrar un equilibrio entre la capacidad de exploración y explotación del algoritmo. La Figura 4.5 muestra los resultados de la optimización.

Figura. 4.5: Resultados del proceso de optimización con el algoritmo JADE con Archivo.



Fuente: Elaboración propia.

El aprendizaje discriminativo alcanzó un valor CLL de -42247.17, en comparación con el -42408.02 obtenido por el aprendizaje generativo. Esta diferencia, aunque marginal, revela que cuando la estructura de la red representa adecuadamente las relaciones de dependencia entre los atributos, ambos enfoques tienden a producir resultados similares. En nuestro caso de estudio, el impacto del Aprendizaje Discriminativo no fue significativamente relevante.

No obstante, hay oportunidades para explorar escenarios donde las relaciones de dependencia entre variables no estén tan bien capturadas en la estructura de la red. En tales situaciones, el enfoque discriminativo podría ser más adecuado para mejorar la generalización del modelo, especialmente si el objetivo principal es analizar e interpretar los resultados de la simulación basada en agentes.

En nuestro estudio, logramos analizar con éxito los resultados del modelo basado en agentes, detallando las interrelaciones entre las variables que representan diversos escenarios de simulación. Esto resalta el verdadero valor del aprendizaje offline a posteriori en contextos de análisis: descifrar y comprender el nexo entre las distintas variables y escenarios. Sin embargo, si la meta es interpretar modelos más opacos o 'cajas negras', existen herramien-



tas especializadas como LIME, Shapley, entre otros, que pueden proporcionar claridad y transparencia en tales casos.

4.3. Apéndice



electronics



Article

Agent-Based Models Assisted by Supervised Learning: A Proposal for Model Specification

Alejandro Platas-López ^{*,†}, Alejandro Guerra-Hernández [†], Marcela Quiroz-Castellanos [†] and Nicandro Cruz-Ramírez [†]

Instituto de Investigaciones en Inteligencia Artificial, Universidad Veracruzana, Xalapa 91000, Mexico

* Correspondence: alejandro.plataslo@anahuac.mx; Tel.: +52-228-221-5879

† Current address: Calle Paseo Lote II, Sección 2a No. 112, Nuevo Xalapa, Veracruz 91097, Mexico.



Citation: Platas-López, A.; Guerra-Hernández, A.; Quiroz-Castellanos, M.; Cruz-Ramírez, N. Agent-Based Modeling Supported by Machine Learning: A Specification Proposal and Case Study. *Electronics* **2023**, *12*, 495. <https://doi.org/10.3390/electronics12030495>

Academic Editors: Anastasia Angelopoulou, Konstantinos Mykoniatis, Sean Mondesire, Athanasios Aris Panagopoulos and Maysam Abbod

Received: 1 December 2022

Revised: 11 January 2023

Accepted: 13 January 2023

Published: 18 January 2023



Copyright: © 2023 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Abstract: Agent-based modeling (ABM) has become popular since it allows a direct representation of heterogeneous individual entities, their decisions, and their interactions, in a given space. With the increase in the amount of data in different domains, an opportunity to support the design, implementation, and analysis of these models, using Machine Learning techniques, has emerged. A vast and diverse literature evidences the interest and benefits of this symbiosis, but also exhibits the inadequacy of current specification standards, such as the Overview, Design concepts and Details (ODD) protocol, to cover such diversity and, in consequence, its lack of use. Given the relevance of standard specifications for the sake of reproducible ABMs, this paper proposes an extension to the ODD Protocol to provide a standardized description of the uses of Machine Learning (ML) in supporting agent-based modeling. The extension is based on categorization, a result of a broad, but integrated, review of the literature, considering the purpose of learning, the moment when the learning process is executed, the components of the model affected by learning, and the algorithms and data used in learning. The proposed extension of the ODD protocol allows orderly and transparent communication of ML workflows in ABM, facilitating its understanding and potential replication in other investigations. The presentation of a full-featured agent-based model of tax evasion illustrates the application of the proposed approach where the adoption of machine learning results in an error statistically significantly lower, with a *p*-value of 0.02 in the Wilcoxon signed-rank test. Furthermore, our analysis provides numerical estimates that reveal the strong impact of the penalty and tax rate on tax evasion. Future work considers other kinds of learning applications, e.g., the calibration of parameters and the analysis of the ABM results.

Keywords: agent-based modeling and simulation; machine learning; ODD protocol; tax evasion

1. Introduction

Modeling is an essential practice in science [1]. A scientific model is an abstract and simplified representation of a phenomenon of interest that makes its main features explicit and visible, and that can be used to understand the phenomenon, develop explanations, and generate predictions. From the natural to the social sciences, scientists have adopted different modeling alternatives. Among them, according to Gilbert [2], the agent-based modeling (ABM) approach is common, since it allows the development of models where individual entities and their interactions are directly represented, and it also makes it possible to model individual heterogeneity, specifying explicitly the rules of behavior of the agents and situating the agents in geographic or other space. This empowers modelers to represent multiple scales of analysis naturally, to ascertain macro- or societal-level structures from individual actions, and to employ various types of adaptation and learning. None of these advantages is easy to replicate with other modeling approaches, offering the modeler a lot of freedom to implement his or her model.

Despite this positive scenario, some specific disadvantages when implementing ABMs have been identified in the literature. They include the complexities involved in its design, the difficulty in calibrating its parameters, the laborious analysis and interpretation of its results, its high computational cost, and the programming skills required for its implementation. With the exception of the last two mentioned, the rest of these topics are related to various stages of ABM development, that is, the design, calibration, experimentation, validation, and analysis stages. Typically, these stages involve data processing before, during, or after the simulation itself. Machine Learning (ML) algorithms can be used in such processing, suggesting that ABM could take advantage of ML, given the greater availability of data, computational resources, and learning algorithms [3].

In recent years, more and more publications have considered the synergy between ML and ABM as a means to overcome some of the mentioned disadvantages. However, these efforts are still somewhat disorganized, as some researchers focus on assisting in model design [4–6], others focus on parameter calibration Chu et al. [7], Lamperti et al. [8], Zhang et al. [9], and still, others aim to find optimal behavior policies [10–12]. So, very often the contributions are isolated and address particular problems rather than the emerging methodological issues that result from the integration of ML and ABM. Therefore, the main objective of this article was to propose a framework that allows the transparent description of the integration of both tools, in addition to providing a broad, but integrated, image of the current state of knowledge on ML–ABM interaction, emphasizing the different applications that exist to address specific problems in the different stages of development of agent-based models.

The rest of the article is structured as follows. Section 2 presents a literature review on the efforts made to synthesize the applications of ML in ABM, as well as computational models on tax evasion. Section 3 provides context by characterizing the criticisms of ABM reported in the literature, and identifying the kinds of ML algorithms considered to overcome these criticisms and the related issues. In Section 3 we also propose an extension to the ODD protocol [13,14] with the aim of overcoming the current limitations for a comprehensive and transparent description of the incorporation of learning in ABM. Section 4 presents a case study on tax evasion of the payroll tax in Mexico, which includes the ODD protocol in the terms proposed to describe the design of an ABM assisted by ML for the inclusion of a synthesized agent from existing data on the Mexican tax system. Section 5 presents the results of the implemented model, in which the versions of the model with and without learning are compared, highlighting the usefulness of synthesizing existing data to better capture the nature of the phenomenon under study. Finally, Section 6 offers conclusions.

2. Literature Review

This section summarizes the work that has been done to characterize the combination of machine learning tools and agent-based modeling. In the same way, since a full-featured agent-based model on tax evasion is presented as a case study, articles that address this phenomenon from the point of view of computer simulation are also presented.

2.1. Machine Learning in Agent-Based Modeling

There are previous works that have analyzed the incorporation of ML in agent-based models from a slightly broader perspective. Müller et al. [15] considered an extension to the ODD protocol, named ODD+D, to include human decision-making. It represents a correct step to complement the description of the agent-based models, which allows transparency and facilitates the understanding of the models. However, it does not refer to the different applications of ML in ABM that can be carried out, and, in any case, it refers only to the design of rules of behavior.

In a first contribution, Elbattah and Molloy [16] focused specifically on ML applications to assist in design, namely, data-driven feedback. However, they did not consider other

types of applications, such as online learning for behavior optimization or post-learning for calibration or analysis of results.

In a second contribution, Elbattah [17] considered three categories for applying ML in agent-based modeling, which they call conceptual modeling, runtime modeling, and output analysis. Although they considered other ways of including ML in ABM, they did not consider a posteriori offline learning for parameter calibration. Also, in neither of their two contributions did they indicate how a researcher could carry out the implementation and documentation of the incorporation of both tools.

Zhang et al. [18] distinguished four scenarios where ML could contribute to the ABM process, two of which were at the agent level (micro), and two at the model level (macro). As for the micro level, they differentiated two applications of ML, the first one being agent situational awareness learning, in which ML is commonly used to predict variables or agent-related behaviors. The second is based on the idea that as an agent lives in a changing environment, agents obtain more knowledge and experiences, which can be stored in the form of data. This data can be used to infer better decision-making models by applying appropriate ML techniques to intervene in the behaviors of this agent, allowing them to better achieve their goals.

At the model level, they mention two applications, the first is useful when the number of agents and parameters grows and the combinations of parameters grow exponentially, which generates a high computational cost, so ML is used to build the mapping structure (emulator) between the parameters and the emergency exit of the ABM. The second is to apply ML techniques to enhance the decision-making of the policymaker at the macro level. In these cases, the macro-level policymaker acts as a single “macro-agent”.

Micro applications map to a priori offline learning and online learning, whereas macro applications map to post-hoc offline learning in their two tasks, validation-calibration, and analysis. However, although all possible tasks are taken into consideration, a methodology that allows researchers to carry out the transparent implementation is not proposed.

2.2. Tax Evasion

Tax evasion is an illegal and intentional activity taken by individuals to reduce their legally due tax obligations [19]. Although formal research on tax evasion began in the 1970s, with the seminal work of Allingham and Sandmo [20], who applied the economics of crime approach [21] to tax reporting behavior, scholars and practitioners continue to face the challenge of designing and implementing policies and incentives to mitigate tax evasion.

Taxes are essential to fund public expenditures. Therefore, tax evasion is not only a problem for the tax authorities but also a problem for society. Given that a country's investment in healthcare, education, national defense, social security, transportation, infrastructure, science, and technology mainly comes from public finances, tax evasion leads to mismatches in public goods [22]. So, the analysis of tax evasion can be used to determine the factors that affect the compliance rate and, ultimately, help the government achieve revenue targets. Daude et al. [23] found evidence of such factors in developing countries. They described two groups of factors: socioeconomic factors, such as age or education, and institutional determinants, like trust in government and satisfaction with the quality of public services.

Although early theoretical research provides a baseline for evaluating assumptions about tax evasion, the basic method of classical formulae, i.e., the use of utility functions and the assumptions of taxpayer homogeneity and rationality, insufficiently describe taxpayer behaviors observed in practice. To overcome this limitation in the study of tax evasion, an alternative, based on agent modeling and simulation, has been applied in this endeavor. So-called agent-based models (ABM) are designed to consider personal preferences and can accommodate a wider variety of internal and external variables to help explore a broader space for compliance results.

Mittone and Patelli [24] created a computational model of tax evasion that considers taxpayers' psychological motives and public sector goods. The authors assumed the exis-

tence of three classes of taxpayers. Each class has a unique utility function that describes its behavior. Taxpayer agents are assigned initially to one of the three categories but may change from one category to another during a model run by imitation of successful behavioral strategies. The model demonstrates that when audits are introduced the additional revenue raised increases the quantity of public goods that can be provided. This information is transmitted to the agents who, in turn, incorporate it into their utility calculations.

Davis et al. [25] analyzed the effect of social norms and enforcement on the dynamics of taxpayer compliance. They developed an analytical and a computational model to evaluate the movement between classes of compliant and non-compliant taxpayers. Their analysis suggested that the effect on compliance of changing enforcement levels depended on whether the taxpayer population was initially compliant or non-compliant. Compliant populations are insensitive to changes in enforcement policies until enforcement becomes sufficiently lax. In contrast, relatively non-compliant populations respond to increased enforcement by gradually increasing compliance. Then, when enforcement becomes sufficiently harsh, there is a sudden shift in equilibrium to very high levels of compliance. After the taxpayer population shifts from compliance to non-compliance, or vice versa, their models predicted that returning to the previous enforcement policy would not cause the population to return to its previous state. Their results also help to explain why taxpayer compliance varies across time and across geographic regions, even under similar enforcement regimes.

Hokamp [26] analyzed tax evasion dynamics and social interactions in heterogeneous populations with an agent-based model. They included aging to incorporate psychological findings with respect to social norms up-dating over a taxpayer's life cycle. The provision of public goods with an exponential utility function was also explored. Their findings supported the view that, when social norm updating is considered, age heterogeneity leads to fluctuations in income tax evasion.

3. Specification of Learning in Agent Based Modeling

Although agent-based modeling is a promising approach among several domains of science, some researchers have encountered drawbacks related to some of the stages of the agent-based modeling cycle. These disadvantages can be summarized as design problems [27–35], calibration–validation [28,29,34], interpretation of results and analysis [31,33–38], high computational cost [35,39–45], as well as necessary programming skills [30,37,46].

To overcome these limitations, some authors incorporated ML techniques at different stages of the development cycle of agent-based models. Broadly speaking, we can differentiate between applications with learning outside of the simulation, called offline learning, and learning during simulation, called online learning. In turn, offline learning can be classified into two aspects: a priori learning, in the ABM design stage, and a posteriori learning, applied in the calibration–validation and analysis phases. In the first case, learning uses existing data to instantiate model components. While in the second case, learning is applied to the data produced by the model, either to calibrate their parameters towards a state of validation of the results or to analyze the results by identifying which inputs produce which outputs and so develop a hypothesis of the simulation behavior. Figure 1 summarizes the different categories of ML applications in Agent-Based modeling. In what follows, some of the ML applications falling into the mentioned categories are briefly described.

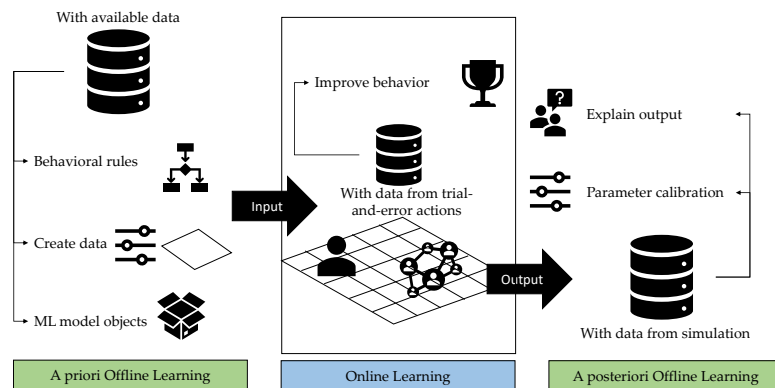


Figure 1. Categories of Machine Learning applications in Agent-Based modeling.

3.1. Design of Agent-Based Models Assisted by Machine Learning

The use of ML for design purposes implies the existence of data to carry out the learning. Among developments found in literature, the following three types of applications are distinguished: the first generates data to later use in the model [4,5,47], e.g., Saeedi [48] develops predictions of land use change with neural networks and, subsequently, implements the results as part of the environment; the second group of applications concerns learning the behavioral rules of the agents, so that the modeler can later implement them in the model [6,49,50]. It is important to note that in this type of implementation, the algorithms that are usually used are easy to interpret and allow the rules to be derived in a simple way, i.e., decision trees or Bayesian networks; and the third group of applications concerns learning an ML model from data, and later inserting that model in the ABM, so that it is available to agents, e.g., [51–53]. Unlike in previous applications, black box models, as artificial neural networks (ANNs), are preferred.

3.2. Validation and Analysis of Agent-Based Models Assisted by Machine Learning

The a posteriori offline learning applications that have been developed use data generated by the simulation. The use of ML for validation purposes consists in finding a subset of input parameters that can produce the desired output. This process is known as parameter calibration, and, in most cases, it has the purpose of performing validation of the model, i.e., ensuring that the model reproduces stylized facts of the real phenomenon. While for analysis purposes, the goal is to understand the complex relationships that occur within the model and to identify which inputs lead to which outputs.

A new approach for calibrating parameters of an agent-based model, using the Bayesian estimation technique to capture the fundamental features of economic fluctuations, is presented in Chu et al. [7]. Using Extreme Gradient Boosting, a tree-based algorithm, which sits under the supervised branch of ML, Lamperti et al. [8] proposes another approach to efficiently calibrating agent-based models to real data. Their proposed method, “learns” a fast surrogate meta-model using a limited number of ABM evaluations and approximates the nonlinear relationship between ABM inputs (initial conditions and parameters) and outputs. A similar approach is presented by Zhang et al. [9] with another ensemble algorithm called CatBoost.

In an attempt to make the problem of parameter tuning more tractable, van der Hoog [54] emulated the input/output function of the entire ABM by a less complex and more tractable Deep Neural Network. This is useful for robustness and parameter sensitivity analysis, since it

allows a much larger exploration of the parameter space. Other approaches involve the use of heuristics [55] and metaheuristics [56] for parameter optimization.

An ABM calibration approach proposed by Ye et al. [57], established a link between agent micro-behavioral parameters and systemic macro-observations. To compress the agent state space, principal component analysis was introduced to avoid the high dimensions of a macrostate transfer equation. Another approach was proposed by Kim et al. [58], in which automatic calibration was done with both dynamic and heterogeneous calibrations, using Hidden Markov Model, Variational Autoencoder, and Gaussian Process Regression.

In order to analyze the spread of *Cutaneous Leishmaniasis* in central Iran, Rajabi et al. [59] proposed an ABM to simulate the dynamics of spread based on a Geographic Automata System, the results of which were analyzed by means of Bayesian modeling. In a similar approach, Hayashi et al. [60] used Decision Trees in order to achieve comprehensive behavior prediction in the case of customer churning in a subscription-based business. With a similar objective, R. Vahdati et al. [61] also used Decision Trees to understand how factors, such as climate, ecology, human behavior, and population dynamics, interacted to affect human survival and dispersal from Africa in the Late Pleistocene age. As can be seen, these approximations use easy-to-interpret or white-box algorithms.

However, the use of black box algorithms is more frequent, e.g., Ozik et al. [62] integrated random forest or artificial neural networks to dynamically and efficiently characterize the parameter space of a large and complex agent-based infectious disease model. Similarly, Chen et al. [63] used random forest regression to evaluate the correlation between parameters and outputs and to then build a metamodel by a neural network to predict the simulation outputs from the parameters. Random forests were also implemented to make inferences from a detailed model representing past human–environment interactions [64]. Garg et al. [65] identified model parameters that were most influential to ABM outputs. Edali and Yücel [66] used the parameters to discover the relationship between inputs and outputs of agent-based models. Gursoy and Badur [67] provided quick predictions on the number of emigrants to understand the effect sizes of individual parameters.

Regarding proposals for analysis through ANN, Xu et al. [68] revealed the learning process and heterogeneity of agents in the urban expansion of Auckland, New Zealand. To predict the time evolution of an ABM, focusing on the clinical condition of sepsis, Xu et al. [68] trained an ANN as a surrogate system. Equally, Xiao and Liu [69] trained an ANN to study the effects of different combinations of control measures for COVID-19. For urban decision-making, Zhang et al. [70] used a convolutional neural network to achieve real-time prediction of an agent-based model. Other proposals considered the use of Support Vector machines and support vector regressions [71], causal discovery [72], reconstructability analysis [73] and the Louvain algorithm [74], an unsupervised algorithm to cluster structures in social networks.

3.3. Online Learning in ABM Assisted by Machine Learning

Online learning makes use of data generated during the simulation. Its purpose is to generate optimal behavior for agents by allowing them to explore the environment through trial and error. Due to the type of learning task, most of the proposals opt for Reinforcement Learning (RL) algorithms. Many works in the literature use Temporal Difference, and, specifically, the Q-learning algorithm [10–12,75–78]. Other Temporal Difference algorithms are also used [79,80]. Liang et al. [81] proposed the Deep Deterministic Policy Gradient Algorithm and Deep Q-Networks were proposed by Sert et al. [82].

A few exceptions adopted supervised learning algorithms. Jäger [83] introduced a framework that, although closely related to RL, used ANN to enable the agents to learn rules. An expectation model of inflation, based on ANN, was introduced by Salle [84], wherein the weights were randomly initialized so that, at least at the beginning, agents' mental models and the resulting expectations differed. As new observations became available, agents' ANN were trained. In the energy market domain, Dehghanpour et al. [85] proposed that each company agent develops a private probabilistic model of the market

using dynamic Bayesian networks. Sparse Bayesian Learning, an online learning algorithm, was used for training the model to infer the future state of the market and to estimate an optimal bidding function.

Although most of the applications have the objective of finding the optimal rule of behavior and using this rule as a hypothesis of the phenomenon, there are some approaches that try to assist the modeler in the design of the ABM by returning the found rule to be implemented. A novel path to agent policy development was examined by Norman et al. [86], in which modelers did not manually craft the policies, but allowed them to emerge through the application of RL within a game engine environment. Fuller et al. [87] proposed the use of RL to enable agents to define their required interaction rules. An approach in which agents created their own optimal strategy and the designer could interpret it was given by Cummings and Crooks [88].

3.4. The Overview, Design Concepts and Details Protocol

The Overview, Design concepts, and Details (ODD) protocol [13,14], is a standardized document providing a consistent, logical, and readable account of the structure and dynamics of ABMs. It consists of seven elements, as shown in Figure 2. These elements are grouped into three levels of description, ranging from general to specific.

Overview. A general description of the model, including its purpose and its basic components, the agents and variables describing them and the environment, and the scales used in the model, e.g., time and space, as well as an overview of the processes and their scheduling.

Design concepts. A brief description of the basic principles underlying the model's design, e.g., rationality, emergence, adaptation, learning, etc.

Details. Full definitions of the involved sub-models.

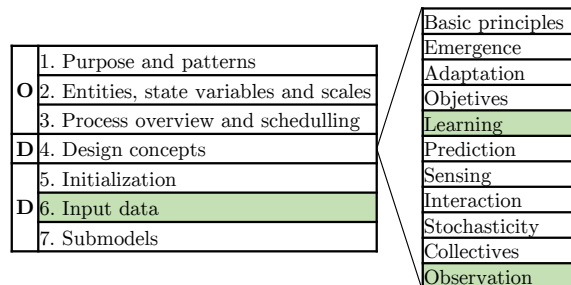


Figure 2. Structure of model descriptions following the ODD protocol. Elements, where an update is proposed, are highlighted. Adapted from [14].

The inclusion of ML in ABM occurred gradually, si the original ODD protocol [89] did not consider learning among the design elements. In the first update to the protocol [13], the relevance of learning was recognized, including it as a design concept, but adopting only the online modality. Instead of enhancing the specification of learning as a design element, the second update to the ODD protocol [14] focused on other aspects, such as clarity, replication, and structural realism. Indeed, in its current state, the only questions about learning in the ODD protocol are the following: “Do many individuals or agents (but also organizations and institutions) change their adaptive traits over time as a consequence of their experience? If so, how?”.

So, the learning section in the current ODD protocol is intended to describe online learning tasks only, leaving out the possibility of documenting offline learning applications,

both before and after simulation, such as those referenced previously. An extension to the ODD protocol enabling the description of all these learning tasks seems necessary.

3.5. Proposal

We propose extending the ODD questions associated with learning as a design element covering:

1. **What is the purpose of learning?** Following the approach of the current ODD Protocol, as stated by Grimm et al. [13], a possible purpose of learning is to improve the performance of the agents during simulation based on their experience. Other equally valid purposes, as discussed above, include the following: generating data for simulations, inducing the rules defining the behavior of the agents from existing data, calibrating parameters of the model, and analyzing the results of the simulations.
2. **When is learning performed?** According to its purpose, learning can take place at different moments. In the current ODD Protocol, learning is expected to occur during the simulation. It is also possible to learn before the simulation, to obtain models to generate synthetic data to be used during simulation, or the rules of the agents included in the ABM. It is also possible for the learning to occur after the simulation to calibrate parameters or to analyze the data produced by the simulation.
3. **What components of the ABM are affected by learning?** In the current ODD Protocol, agents are the only components performing, and being affected by, learning. However, it is also possible that the environment exploits the use of ML methods, e.g., consuming synthetic data generated by learned models. It is also possible to learn the values of the parameters of the ABM to obtain some desired behavior of the whole system.
4. **How is learning computed?** Finally, the type of algorithm used and its input data must be considered. Online learning requires incremental learning algorithms, such as those included in, but not limited to, RL. Supervised techniques can be used for offline learning, both before and after simulation. Evolutionary approaches are well suited for calibrating parameters. The preference for explainable methods versus black box techniques is also related to the purpose of learning.

Some considerations must be addressed in other sections of the ODD protocol to obtain a more comprehensive description of the learning process. They are described as follows.

A priori ML implies the use of existing data related to the phenomenon to be simulated. The use of such data must be included in the details section of the ODD Protocol, specifically in the element called “Input data”, in which the following question is answered:

6 Input data

Current: Does the model use input from external sources, such as data files or other models, to represent processes that change over time?

Proposal: Does the model use data from external sources (data files, or other models) to represent **some element in the model**? **Is the data processed with machine learning**?

A posteriori ML implies the use of data resulting from the simulation. The use of such data must be included in the design concepts section of the ODD Protocol, and, specifically, in the element called “Observation”, in which the following question is answered:

5.11 Observation

Current: What data are collected from the ABM for testing, understanding, and analyzing it, and how and when are they collected? Are all output data freely used, or are only certain data sampled and used, to imitate what can be observed in an empirical study?

Proposal: What data are collected from the ABM for **calibrating**, testing, understanding, and analyzing it, and how and when are they collected? Are all output data freely used, or are only certain data sampled and used, to imitate what can be observed in an empirical study? **Is the data processed with machine learning?**

The questions rewritten in this way link the source of the data with the learning section in which the learning task to be performed is described. Finally, a complete description of the ML model must be provided in the ODD protocol sub-models, including data preparation, modeling, and evaluation.

4. Case Study: The Payroll Tax Evasion Model

This section describes in detail an ABM in which the design is assisted by ML. This specification adheres to the extension proposal to the ODD protocol made in the previous section. It is important to highlight that the case study is a full-featured model, in which the learning application has favorable repercussions for the analysis of fiscal policy. As it is an application of ABM design, assisted by ML, the aspects of the ODD where the improvements are made are in the learning design concept, the input data, and the sub-model, Section 2, where the ML algorithm used is described in detail.

The purpose of this model is to advance the state of the art in tax evasion analysis, presenting an alternative based on agent simulation and machine learning. Where certain changes in taxpayer behavior can be mined from data and applied during the simulation, which otherwise would not be detected by the model. Evasion of payroll tax in Mexico is modeled as a case study to exemplify the benefits of the combination of ML and ABMs. Each Mexican state has autonomy over the way in which payroll tax is collected. Therefore, to model these different fiscal scenarios and their effects, an explicit representation of the space is made through a Geographic Information System with hexagonal tessellation.

For the sake of reproducibility, the model is described in the following sections using the ODD protocol [13,14].

4.1. Overview

Purpose. Analyze the effect of institutional determinants on payroll tax evasion in Mexico.

Entities, state variables, and scales.

- Agents: Employers and tax authority.
- Environment: A Mexican state-level representation of the payroll tax system.
- Scales: Time is represented in discrete periods, each step representing a month, which corresponds to the tax collection period according to the current legislation. In order to have a margin of error less than 5% with a confidence interval of 99% in the selected sample, each employee agent in the model represents 2000 employers in the 2019 Mexican labor market.
- State variables: The attributes that characterize each agent are shown in Table 1 along with the method for initializing each variable.

Process overview and scheduling.

1. Employers decide in which market to, formally or informally, query a machine learning model.

2. Informal employers are full evaders. Formal employers calculate the amount of taxes to report. If they calculate so as to report zero taxes, they also become full evaders. If they calculate to report all the tax, they become full taxpayers. In any other case, they become partial evaders.
3. Tax authority collects the declared amount of taxes.
4. The tax authority conducts audits on a random basis. If the audit is successful, then partial and full evaders must pay the evaded amount and a penalty for the undeclared amount.
5. Every 12 months employers increase their age. With some probability, in each period, employers can die. If this happens, they are replaced by another employer with the same characteristics, except for age.

4.2. Design Concepts

Basic principles. Voluntarily declared income rises with increasing individual income, penalty rate, and audit probability, and decreases with increasing tax rate, as proposed in the neoclassical theory of evasion by Allingham and Sandmo [20], and also includes exponential utility and public goods provision, described by Hokamp [26]. The distribution of production among employers follows a power law in which there is a marked inequality in the distribution, characterized by a small percentage of people holding the most resources [90]. This characterizes a capitalist economy, such as that in Mexico. Mortality of employers follows a Weibull distribution [91]. This function is adjusted for the case of Mexico, where life expectancy at birth is seventy-five years [92].

Emergence. The extension of tax evasion is an emerging result of the adaptive behavior of employers. These results are expected to vary in complex ways when the characteristics of individuals or their environment change. Other outcomes, such as the distribution of production and age, are more strictly imposed by rules and, therefore, less dependent on what individuals do. However, these results are important for the validation of the model.

Adaptation. Employers can change their alignment in the formal sector. To make that decision, they query a previously trained machine learning model. With this trait, agents could become full evaders, and, therefore, avoid paying taxes. The tax authority does not have adaptation mechanisms.

Objectives. The employer's goal is to increase his utility by paying the least taxes. The objective of the tax authority is to increase tax collection.

Learning. With the extension proposal for learning inclusion, the following questions are answered. *What is the purpose of learning?* The purpose of learning is to deduce the rules defining the behavior of the agents from existing data. *When is learning performed?* Learning is performed before the simulation to obtain models for generating the rules of the agents included in the ABM. *What components of the ABM are affected by learning?* The environment is affected by encapsulating, in a learned model, the data of the fiscal system and the quality of public goods. *How is learning computed?* Learning is computed by a supervised algorithm.

Prediction. The employers' learning is an offline process. During the simulation, employers sense their current conditions and act accordingly. Employers do not have an internal model to estimate their future conditions or the consequences of their decisions.

Sensing. Employers consider the size of their company, education, sector of occupation, state, size of the locality, age, and income as internal states, and they perceive tax rate, and level of insecurity as environmental variables. If an audit is successful, the tax authority can collect undeclared tax from employers.

Interaction. Interactions between employers and the tax authority are direct when, with a certain probability, an audit is carried out.

Stochasticity. The audit process is assumed to be random, with both the probability of it being carried out and its probability of success following a uniform distribution. The process of initialization of the employers' income is also considered random, following a power law. The probability of death of employers is also random following a Weibull

survival distribution. In the first process, stochasticity is used to make events occur with a specific frequency. In the last two processes, stochasticity is used to reproduce the variability in processes for which it is not important to model the real causes of the variability.

Collectives. Employers are grouped into three types of tax behavior: full evaders, partial evaders, and full taxpayers. This collective is a definition of the model in which the set of employers with certain properties about the amount of taxes paid is defined as a separate kind of employer with its own variables and traits.

Observation. At the end of each run, data is collected on the extent of tax evasion, the amount of tax collected, and the number of full evader employers.

4.3. Details

Initialization. At the time $t = 0$, of every simulation run, $N = 1337$ employers were initialized and distributed in each state according to the input data. A total of 32 auditors were also initialized, representing each state tax authority. Some state variables were initialized by a submodel, either deterministic or random, while others were based on data, as shown in Table 1.

Table 1. State variables and method for initialization. In parentheses, is the name of the attribute in the National Institute of Statistics and Geography (INEGI) dataset [93] of state variables.

Agent	Attributes	Type	Initialization	Value
Auditor	penalty-collected	Float	Deterministic	0
	tax-collected	Float		0
	my-employers	AgSet	Random	Submodel 5
	ent-auditor	Int		[1, 32]
Employer	business-size (ambito2)	Int	Database	{0, 2, 3, 4, 5, 8}
	education (anios_esc)	Int		[0, 20]
	economic-activity (c_ocu11c)	Int		[1, 10]
	age (eda)	Int		[17, 98]
	mexican-state (ent)	Int		[1, 32]
	income (ing7c)	Int		[1, 7]
	formal-or-informal (mh_col)	Int		[0, 1]
	size-of-region (t_loc)	Int		[1, 4]
	corruption	Float		(0, 1)
	insecurity	Float		(0, 1)
	tax	Float		(0, 1)
	audit?	Bool	Deterministic	false
	audited?	Bool		false
	type-of-taxpayer	Int		2
	declared-tax	Float		0
	payroll	Float	Submodel 11	Submodel 11
	payroll *	Float		Submodel 12
	risk-aversion- ρ	Float		Submodel 16
	undeclared-payroll	Float		0
	undeclared-tax	Float	Random	0
	α -s	Float		0.05
	δ	Float		−0.1
	prob-formal	Float		(0, 1)
	production	Float		Submodel 10

In the same way, input parameters of simulation were adopted from literature, databases or through experimentation. Table 2 shows the initial values of the baseline model. In Mexico, the penalty rate varies in some states between 55 and 75%, and in others between 75 and 100%, so an initial value of 75% was adopted. The variation in tax rates, perception of insecurity, and corruption started at zero, that is, the value indicated in local tax legislation and in INEGI was taken. According to Bonet and Rueda [94], the effectiveness of tax collection in Mexico is approximately 0.7. For the decision threshold, 0.5 was used, since, according to [95], the threshold of 0.5 is commonly used to estimate the

probability for a set with two classes. For the audit probability and the effectiveness of the audits, values were found through experimentation to guarantee outputs similar to the dynamics of the Mexican tax system, which is small but efficient. The user interface also had other input parameters, one of them allowing switching the machine learning model “on” and “off”.

Table 2. Input parameter initialization of baseline model.

Parameter	Description	Value	Initialization
π	Penalty rate	0.75	Database
α	Audit probability	0.05	Experimentation
ϵ_{AP}	Effectiveness of audit process	0.75	Experimentation
ϵ_{TC}	Effectiveness of tax collection	0.70	Literature [94]
$\Delta\theta$	Variation in tax rate	0.00	Database
ΔPI	Variation in perceived insecurity	0.00	Database
ΔPC	Variation in perceived corruption	0.00	Database
τ	Threshold for formal or informal sector choice	0.50	Literature [95]

Input data. Does the model use data from external sources (data files, or other models) to represent some element in the model? Yes, the model used external input data to learn the Mexican tax system, that is, the National Survey of Occupation and Employment (ENOE), and the National Survey of Quality and Government Impact (ENCIG), as well as the tax laws of the different states where the payroll tax rate was specified. The period considered for the 3 sources of information was from 2011 to 2019.

ENOE is a quarterly survey, but just the third quarter was used as a reference for annual information. Among other data, it offers sociodemographic variables on the characteristics of the employers. From these variables, it could be determined whether an employer was in the formal or informal sector. ENCIG is a biannual survey in which people are asked about the top 3 (three) problems they believe exist in their state. The main problems among all the states were insecurity, corruption, and unemployment. Insecurity and corruption were summarized by state and interpolated to get annual information. The resulting data was joined to the ENOE data set along with the tax data. The dataset collected in this way gave a matrix of size $71,833 \times 10$, where the proportion of formal employers was 60.57%.

Is the data processed with machine learning? Yes, from the resulting pre-processed ENOE data, sampling was performed using the local pivot method with the algorithm offered by the R package “SamplingBigData” [96], which effectively generated a balanced sample data set. Selected attributes to generate the balanced sample were the state, the employer’s classification in the formal or informal sector, and the size of the employer’s economic unit. This ensured that the employers in the model reflected the actual proportions of the Mexican labor market. The size of the selected sample corresponded to a scale of 1 to 2000. Figure 3 shows the ML integration process.

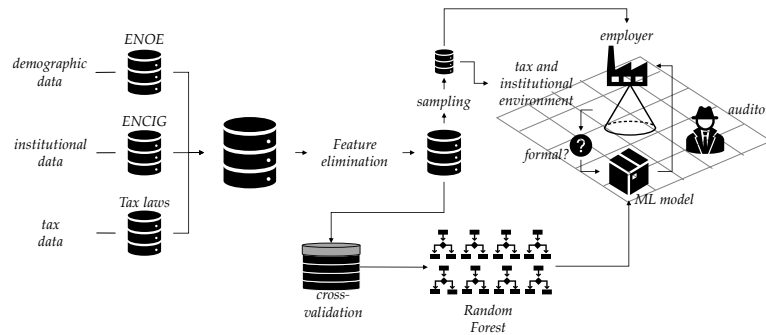


Figure 3. Diagram of integration of ML in the Payroll Tax Evasion Model.

Submodels.

1. After performing a pre-processing of the database, including a manual selection and a recursive feature elimination with resampling algorithm [97] available in R package “caret” [98], it was determined that the main variables that determined the employer’s sector were the size of the business (ambito2), education (anios_esc), economic activity (c_ocu11c), state (ent), size of the region (t_loc), and age (eda).
2. A fast implementation of Random Forest in the R package “ranger” [99] was chosen to learn from data because it provides fast model fitting and evaluation, is robust to outliers, can deal with simple linear and complicated nonlinear associations, and produces competitive prediction accuracy [100]. To tune the hyperparameters and evaluate the performance of the model, cross-validation with $k = 10$ folds was carried out. The final hyperparameter values were $mtry = 20$, $ntrees = 100$, and $nodesize = 1$. That setting provided an accuracy of 83.79%, which was considered good to avoid overfitting. The trained model was available to employers during the simulation.
3. A Geographical Information System layer was loaded. Each polygon was a hexagonal tessellation of the corresponding Mexican state.
4. $N = 1337$ employers were generated and initialized with information from the database and moved to their corresponding state.
5. Auditors were generated and located in their assigned state.
6. The a priori learned random forest model was loaded.
7. Pareto-law values with the following distribution function were generated:

$$f(x) \sim x^{-1-\gamma}$$

8. Where γ was known as the Pareto exponent and estimated to be $\approx 3/2$ to characterize a capitalist economy [90].
9. x are the values generated by a normal distribution function with a mean of 2 and a standard deviation of 0.2 for informal employers and 0.3 for formal ones.
10. To assign a fixed monthly production value to each employer. Generated power law values were multiplied by 23 in the case of informal employers and 50 for the formal. Those quantities generated a perfectly mixed Pareto distribution according to the basic principles and preserved the participation of the informal economy in Mexican Gross Domestic Product (GDP) [101].
11. For simplicity, it was assumed that each employer allocated 30 percent of the value of production to payroll W . The share of wages in Mexican GDP was between 30 and 40% [102].
12. At the beginning of the simulation, it was assumed that non-informal employers declared all the tax, i.e., declared payroll $W^* = W$.

13. At the beginning declared tax X^* by each employer was equal to the declared payroll W^* multiplied by the tax rate θ in the employer's state.
14. Every 12 periods (months) employers increased their age, and they consulted the learned model to decide whether to opt for the formal or informal market, taking their internal attributes and perceived insecurity as a reference.
15. Informal employers did not declare taxes.
16. By social norm [26], employers modified their risk aversion ρ according to their age, as follows:

$$\rho \sim \begin{cases} U(0.0, 0.25) & \text{if } \text{age} \leq 34 \\ U(0.25, 0.5) & \text{if } 34 < \text{age} \leq 51 \\ U(0.5, 0.75) & \text{if } 51 < \text{age} \leq 67 \\ U(0.75, 1.0) & \text{if } \text{age} \geq 67 \end{cases}$$

17. Let β the perceived public goods efficiency, and π the penalty rate.
18. Let ϵ_{AP} and ϵ_{TC} the effectiveness of audit process and tax collection, respectively.
19. Let α the true audit probability and α_S the subjective audit probability known to the employer.
20. Let $\delta = 0.1$, the updating parameter for α_S .
21. If an employer was audited in a specific period, subjective audit probability became 1.
22. In each period (if not audited again) α_S decreased in δ amount until $\alpha_S = \alpha$.
23. In each period, employers calculated the amount of taxes to declare voluntarily X^* , applying the expected utility maximization procedure adopted by Allingham and Sandmo [20]. Let lower bound be:

$$\alpha_S > \frac{1}{1 + \left(\frac{(1-\beta(1-\epsilon_{AP}))\pi}{(1-\beta(1-\epsilon_{TC}))\theta} - 1 \right) e^{\rho(1-\beta(1-\epsilon_{AP}))(\pi W)}}$$

24. And the upper bound be:

$$\alpha_S < \frac{1}{1 + \left(\frac{(1-\beta(1-\epsilon_{AP}))\pi}{(1-\beta(1-\epsilon_{TC}))\theta} - 1 \right) e^{\rho(1-\beta(1-\epsilon_{AP}))(\pi W)}}$$

25. If the subjective audit probability α_S exceeded the upper limit in submodel 22, the employer became fully tax compliant, that is, $X^* = W\theta$, and when α_S fell below the lower bound in submodel 21, the employer fully evaded, that is $X^* = 0$.
26. For α_S in the range for an inner solution, the employer voluntarily declared:

$$X^* = W - \frac{\ln \left(\frac{(1-\alpha_S)(1-\beta(1-\epsilon_{TC}))\theta}{\alpha_S((1-\beta(1-\epsilon_{AP}))\pi - (1-\beta(1-\epsilon_{TC}))\theta)} \right)}{\rho\pi(1-\beta(1-\epsilon_{AP}))}$$

27. The tax authority collected payroll taxes that employers voluntarily declared.
28. The tax authority carried out audits with a random probability of α and a level of effectiveness ϵ_{AP} .
29. If an evader was detected the undeclared tax was collected and a penalty rate π applied over the undeclared tax.
30. In each period, employers had a probability of dying, following a Weibull quantile derivation function:

$$Q(p) = \lambda \left[\frac{1}{1-p} \right]$$

31. Where $\lambda = 0.019$ and $k = 0.479$ are the scale and shape parameters, respectively.

32. It was assumed that, when an employer died, someone else took their place with the same attributes, except for age, which was generated according to:

$$eda = \lfloor X \rfloor$$

$$X \sim N(\mu, \sigma^2) \sim N(37, 6)$$

33. At each time t , the observed output Extent of Tax Evasion (ETE) was calculated as follows:

$$ETE_t = 1 - \frac{\sum_{i=1}^N W^*}{\sum_{i=1}^N W}$$

5. Results

The model was implemented in Netlogo [103]. Figure 4 shows part of the resulting GUI that allowed the initialization of global parameters and provided a view of the agents in a grid environment. The color of the hexagons represented the tax collected by the entity. Employing agents were colored blue, red, or cyan, depending on whether they were tax compliant, evading, or partially evading, respectively. The monitors on the right side show the outputs produced in the model with the given parameter settings.

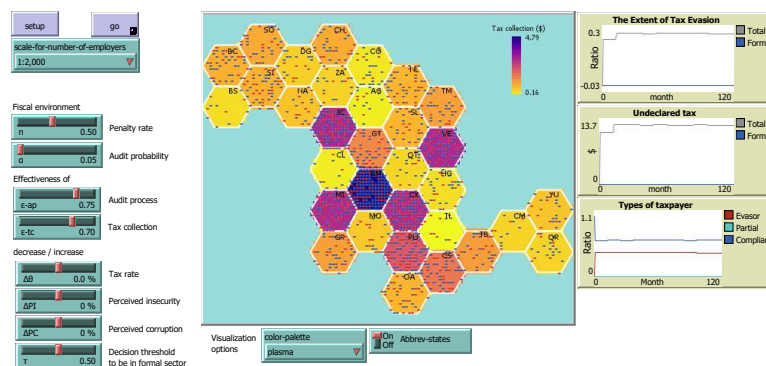


Figure 4. Graphical user interface of the model.

Although the demonstrated simulation model is useful, for “what-if” studies that explore different scenarios, policy analysis requires validated descriptive simulation models [104]. In this section, a validation of the results of the baseline model was carried out in micro and macro terms. Validation consists of replicating the statistical properties of the observed data. After validation, a comparison was performed between the simulation outputs of the models with and without machine learning. Assessment was made in terms of the extent of tax evasion, collected taxes, and the proportion of employers that fully evaded.

5.1. Validation

Empirical validation of agent-based models in economics has undergone significant development [105]. Among the different validation proposals, the indirect approach is one of the most used and the first evaluation that an ABM must satisfy. The core of indirect validation is that micro and macro variables can generate stylized facts or statistical properties that the modeler can compare with those obtained from empirical analysis of the corresponding real-world data set [105,106]. An indirect approach was used to carry out the validation since there were output variables that should follow stylized facts, as well as real data to carry out the comparison.

At the micro level, validation consists of verifying the existence of some stylized facts in economics. As mentioned previously, one of the basic principles lies in the assumption that, in a capitalist economy, the distribution of production follows a power law. This distribution is also transferred to the distribution of the payroll, assuming that greater production demands a greater amount of wages. Figure 5 shows that submodels 5 to 9, which adhered to the basic principle of the Pareto distribution, generated a power law in the payroll distribution, i.e., a positive bias in the distribution. Figure 6 shows the positive bias of 25 independent runs, which was measured with the skewness, a measure of the asymmetry of the probability distribution of a real-valued random variable about its mean [107].

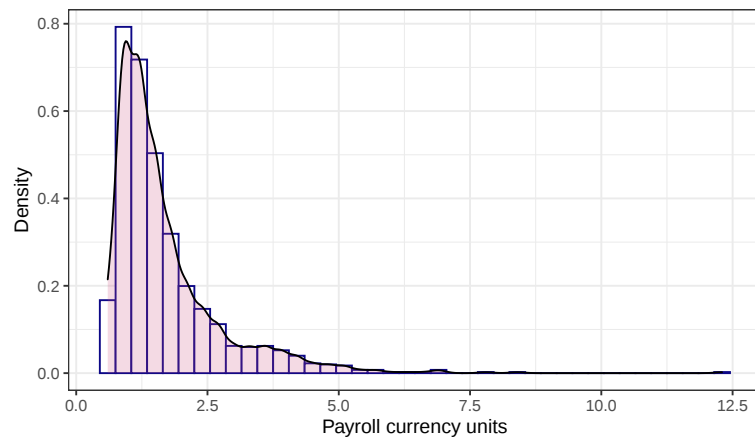


Figure 5. Distribution of payroll follows a power law.

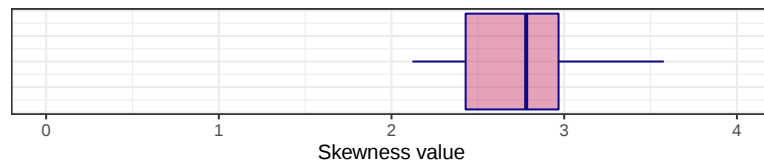


Figure 6. Skewness values obtained over 25 independent runs of payroll distribution. It was considered that values greater than 1 corresponded to highly positively skewed distributions.

At the macro level, a comparison was made between the amounts of tax collection obtained in the simulation and those reflected in the official statistics by state. The data on tax collection that had simulation as a source, were divided into simulations with machine learning “on” and “off”. For the simulation data, the values corresponded to the arithmetic mean of 25 independent runs of the values simulated last year. The values of the real source corresponded to those reported by INEGI in 2019 [108]. Since the “real” data source contained quantities much larger than those produced in the model, all the data were centered and scaled to make the data comparable. Figure 7 shows the distribution of payroll tax collection by the data source. Visually, the distribution of the payroll tax collection approximated the data found in the statistics.

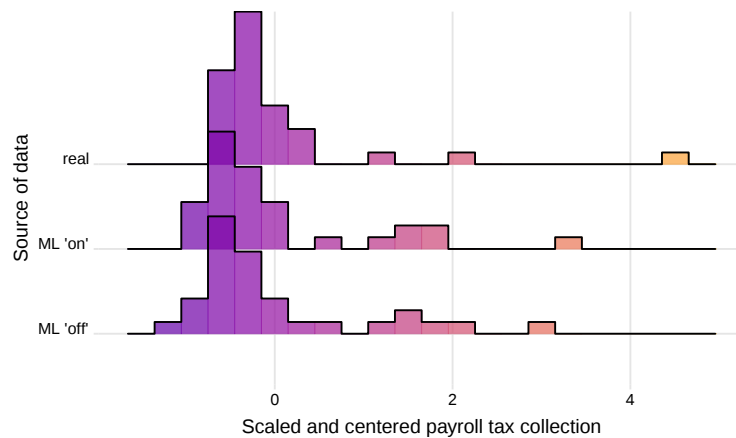


Figure 7. Distribution of scaled and centered payroll tax collection by the source of data.

The values of the state collection of the payroll tax from the statistics for the period 2010 to 2019, and from the arithmetic mean of 25 independent runs of the simulated last year with and without machine learning, were ranked from least to highest. The monotonic relationship between these data sources was measured with Spearman's correlation coefficient. In politics, absolute values greater than, or equal to, 0.6 are considered a strong correlation [109]. Table 3 shows the Spearman correlation coefficients between these data. A strong correlation was obtained for all the years considered. Both micro and macro validations demonstrated that the baseline model was a promising tool for analyzing payroll tax evasion.

Table 3. Input parameter initialization of baseline model.

Model	Year of Real Source									
	2010	2011	2012	2013	2014	2015	2016	2017	2018	2019
ML 'off'	0.73	0.75	0.68	0.68	0.64	0.61	0.61	0.60	0.60	0.61
ML 'on'	0.71	0.73	0.68	0.68	0.65	0.62	0.60	0.60	0.60	0.61

Finally, the Root Mean Squared Error (RMSE) was computed to validate the results of the models with and without machine learning, when compared to the actual values. Figure 8 summarizes by Mexican state and model, the error between the predicted annual tax collection in 10 simulated years and the actual collection in the period 2010 to 2019. The RMSE in the model with machine learning was lower in 19 of the 32 states. To determine if these differences were significant, a test was performed to compare a location measure, since the error of both samples, for the treatments without and with learning, were not normal. The Wilcoxon signed-rank test was used. The alternative hypothesis was that the median RMSE without learning was greater than the median RMSE after learning. The p -value of the test was 0.02166, which was less than the α significance level of 0.05. We could conclude that the median RMSE before ML was significantly higher than the median RMSE after ML.

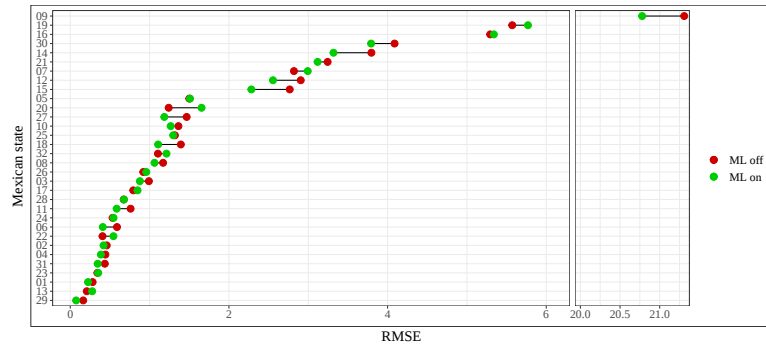


Figure 8. Root Mean Square Error (RMSE) of predicted and actual payroll tax collection in 10 simulated years by state with machine learning “on” and “off”. The x-axis break was done to effectively utilize plotting space and deal with outliers [110].

5.2. The Effect of Machine Learning in Simulation

Figure 9 shows the variations in the outputs of the three observed variables when ML was turned “on” or “off” and changes in the perceived corruption or insecurity were introduced. The value observed in the data corresponded to the coordinate 0.0. Negative values represented a decrease in the perceived value in the respective institutional factor, i.e., less corruption or insecurity. As discussed above, based on the evidence found in developing countries, institutional determinants, such as trust in government and quality of public services, are expected to influence tax evasion [23].

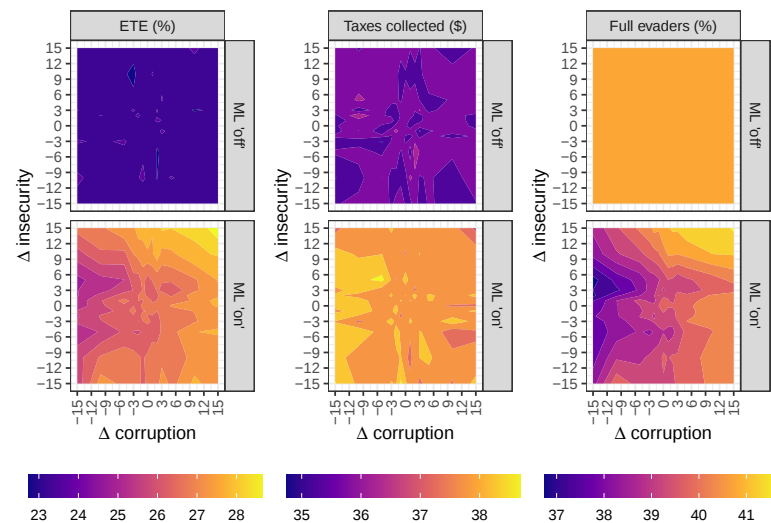


Figure 9. Percentage of ETE, monetary units of taxes collected, and percentage of full evaders in the system, when varying perception on corruption and insecurity, with and without machine learning. Purplish regions denote a small value of the output variable. A low value of ETE or full evaders was good, while, for taxes collected, a low value was bad.

However, it was appreciated that the results of the simulation with ML “off” failed to capture the changes in institutional determinants, mainly in two of the observed variables, the Extent of Tax Evasion, and the number of full evaders in the system. This meant that the analytical model gave greater importance to fiscal parameters (tax rate and penalty rate), and less to institutional parameters (variations in perceived corruption and insecurity).

Conversely, the model with ML “on” could reflect the variations of the institutional factors in a plausible way. For instance, if both the perception of corruption and insecurity increased, the Extension of Tax Evasion also increased. Paradoxically, at that point, there was also a high collection, generated by the greater number of penalties collected. It could also be noted that the perception of corruption played an important role in determining whether employers aligned themselves in the formal or informal sectors. Few evaders in the system and high collection were achieved if the perception of corruption decreased, even if the perception of insecurity increased slightly.

6. Discussion

Different applications of ML in ABM have been reported in the literature. We categorized them according to the moment learning was performed, resulting in online applications, i.e., learning during the simulation, and offline applications, i.e., learning before (a priori) or after (a posteriori) the simulation. Each case had different requirements. A priori learning used existing data about the application domain and might be performed before the ABM even existed to assist its design. A posteriori learning used data generated by the simulation to calibrate parameters and/or analyze the obtained results. On-line learning used data collected during the simulation to enhance the performance of the agents in a given task. For this, different kinds of supervised, evolutionary, and reinforcement learning algorithms were successfully applied.

Reproducible ABMs result from understandable and complete published descriptions. Standardized descriptions, such as those provided by the adoption of the ODD Protocol, contribute greatly in this sense. However, in its current state, the ODD Protocol can not accommodate the diversity of applications of ML in ABMs. Fortunately, the ODD Protocol covers other necessary elements to carry out such descriptions, simplifying its extension with the goal of documenting the diverse applications of ML identified in the ABM literature. For this, the proposed categorization of the ML applications in ABM was central, enabling the design questions guiding the extension to the learning component of the protocol, i.e., What is the purpose of learning? When is learning performed? What components of the ABM are affected by learning? How is learning computed?

The benefits of this extension to the ODD Protocol were illustrated with a case study. A full-featured ABM of the institutional factors on payroll tax evasion in Mexico. In this model, supervised ML was used to synthesize data from the tax environment and encapsulated itself as an agent that would be available to the rest of the agents in the model to make decisions about remaining in the formal or informal sector. The proposed extensions to the ODD protocol made it possible to communicate the ML workflow in the ABM in an orderly and transparent way, favoring its understanding and potential replication in other investigations.

Finally, analyzing the case study, it is important to highlight the contribution of ML to the understanding of the phenomenon. Payroll Tax in Mexico has a low rate, but a high penalty for amounts evaded. With these characteristics, analytical models fail to capture the differences produced by institutional determinants, which seem to have less importance in the mathematical formulation, but have been shown to, in fact, have an effect. The inclusion of ML in ABM allows agents to make decisions that better approximate the complete dynamics of the systems. Therefore, if the intention was to influence the decision of workers and help reduce the extent of tax evasion, a policy recommendation would be related to the enhancement of institutions and the improvement of the quality of public goods provided.

Future work includes applying the proposed extension of the ODD protocol to study cases considering other kinds of learning applications, e.g., the calibration of parameters and the analysis of the ABM results. With the objective of summarizing the basic concepts of Machine Learning and its applications in agent-based modeling in recent years, an analytical and systematic review is also contemplated. This would allow us to determine the prospects and challenges for agent-based modeling assisted by machine learning in the near future.

Author Contributions: Conceptualization, A.P.-L., A.G.-H., M.Q.-C. and N.C.-R.; Investigation, A.P.-L.; Writing—original draft, A.P.-L. and A.G.-H.; Writing—review & editing, A.P.-L., A.G.-H., M.Q.-C. and N.C.-R. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Informed Consent Statement: Not applicable.

Data Availability Statement: Not applicable.

Acknowledgments: The first author thanks the National Council of Science and Technology (CONACYT, Mexico) for the scholarship 743662.

Conflicts of Interest: The authors declare no conflict of interest.

Abbreviations

The following abbreviations are used in this manuscript:

ML	Machine Learning
ABM	Agent-Based Model
ODD	Overview, Design concepts, Details
RL	Reinforcement Learning
ANN	Artificial Neural Network
GDP	Gross Domestic Product
ETE	Extent of Tax Evasion

References

1. Schwarz, C.V.; Reiser, B.J.; Davis, E.A.; Kenyon, L.; Achér, A.; Fortus, D.; Shwartz, Y.; Hug, B.; Krajcik, J. Developing a learning progression for scientific modeling: Making scientific modeling accessible and meaningful for learners. *J. Res. Sci. Teach.* **2009**, *46*, 632–654. [\[CrossRef\]](#)
2. Gilbert, N. *Agent-Based Models*; SAGE Publications, Inc.: Newbury Park, CA, USA, 2020. [\[CrossRef\]](#)
3. Russell, S.; Russell, S.; Norvig, P. *Artificial Intelligence: A Modern Approach*; Pearson series in artificial intelligence; Pearson: New York, NY, USA, 2020.
4. Ghaffarian, S.; Roy, D.; Filatova, T.; Kerle, N. Agent-based modelling of post-disaster recovery with remote sensing data. *Int. J. Disaster Risk Reduct.* **2021**, *60*, 102285. [\[CrossRef\]](#)
5. Yao, F.; Zhu, J.; Yu, J.; Chen, C.; Chen, X. Hybrid operations of human driving vehicles and automated vehicles with data-driven agent-based simulation. *Transp. Res. Part D Transp. Environ.* **2020**, *86*, 102469. [\[CrossRef\]](#)
6. Augustijn, E.W.; Abdulkareem, S.; Sadiq, M.; Albabawat, A. Machine Learning to Derive Complex Behaviour in Agent-Based Modelling. In Proceedings of the 2020 International Conference on Computer Science and Software Engineering (CSASE), Duhok, Iraq, 16–18 April 2020; Institute of Electrical and Electronics Engineers Inc.: Piscataway, NJ, USA, 2020; pp. 284–289. [\[CrossRef\]](#)
7. Chu, Z.; Yang, B.; Ha, C.; Ahn, K. Modeling GDP fluctuations with agent-based model. *Phys. A Stat. Mech. Its Appl.* **2018**, *503*, 572–581. [\[CrossRef\]](#)
8. Lamperti, F.; Roventini, A.; Sani, A. Agent-based model calibration using machine learning surrogates. *J. Econ. Dyn. Control* **2018**, *90*, 366–389. [\[CrossRef\]](#)
9. Zhang, Y.; Li, Z.; Zhang, Y. Validation and Calibration of an Agent-Based Model: A Surrogate Approach. *Discret. Dyn. Nat. Soc.* **2020**, *2020*, 6946370. [\[CrossRef\]](#)
10. Lakić, E.; Artač, G.; Gubina, A. Agent-based modeling of the demand-side system reserve provision. *Electr. Power Syst. Res.* **2015**, *124*, 85–91. [\[CrossRef\]](#)
11. Jalalimanesh, A.; Shahabi Haghighi, H.; Ahmadi, A.; Soltani, M. Simulation-based optimization of radiotherapy: Agent-based modeling and reinforcement learning. *Math. Comput. Simul.* **2017**, *133*, 235–248. [\[CrossRef\]](#)

12. Jalalimanesh, A.; Haghighi, H.; Ahmadi, A.; Hejazian, H.; Soltani, M. Multi-objective optimization of radiotherapy: Distributed Q-learning and agent-based simulation. *J. Exp. Theor. Artif. Intell.* **2017**, *29*, 1071–1086. [\[CrossRef\]](#)
13. Grimm, V.; Berger, U.; DeAngelis, D.L.; Polhill, J.G.; Giske, J.; Railsback, S.F. The ODD protocol: A review and first update. *Ecol. Model.* **2010**, *221*, 2760–2768. [\[CrossRef\]](#)
14. Grimm, V.; Railsback, S.F.; Vincenot, C.E.; Berger, U.; Gallagher, C.; DeAngelis, D.L.; Edmonds, B.; Ge, J.; Giske, J.; Groeneveld, J.; et al. The ODD Protocol for Describing Agent-Based and Other Simulation Models: A Second Update to Improve Clarity, Replication, and Structural Realism. *J. Artif. Soc. Soc. Simul.* **2020**, *23*, 7. [\[CrossRef\]](#)
15. Müller, B.; Bohn, F.; Drefßler, G.; Groeneveld, J.; Klassert, C.; Martin, R.; Schlüter, M.; Schulze, J.; Weise, H.; Schwarz, N. Describing human decisions in agent-based models – ODD + D, an extension of the ODD protocol. *Environ. Model. Softw.* **2013**, *48*, 37–48. [\[CrossRef\]](#)
16. Elbattah, M.; Molloy, O. ML-Aided Simulation. In Proceedings of the 2018 ACM SIGSIM Conference on Principles of Advanced Discrete Simulation, Rome, Italy, 23–25 May 2018; ACM: New York, NY, USA, 2018. [\[CrossRef\]](#)
17. Elbattah, M. How can Machine Learning Support the Practice of Modeling and Simulation?—A Review and Directions for Future Research. In Proceedings of the 2019 IEEE/ACM 23rd International Symposium on Distributed Simulation and Real Time Applications (DS-RT), Cosenza, Italy, 7–9 October 2019. [\[CrossRef\]](#)
18. Zhang, W.; Valencia, A.; Chang, N.B. Synergistic Integration Between Machine Learning and Agent-Based Modeling: A Multidisciplinary Review. *IEEE Trans. Neural Netw. Learn. Syst.* **2021**, 1–21. [\[CrossRef\]](#)
19. Alm, J. Measuring, explaining, and controlling tax evasion: Lessons from theory, experiments, and field studies. *Int. Tax Public Financ.* **2011**, *19*, 54–77. [\[CrossRef\]](#)
20. Allingham, M.G.; Sandmo, A. Income tax evasion: A theoretical analysis. *J. Public Econ.* **1972**, *1*, 323–338. [\[CrossRef\]](#)
21. Becker, G.S. Crime and Punishment: An Economic Approach. *J. Political Econ.* **1968**, *76*, 169–217. [\[CrossRef\]](#)
22. Slemrod, J.; Yitzhaki, S. Tax Avoidance, Evasion, and Administration. In *Handbook of Public Economics*; Elsevier: Amsterdam, The Netherlands, 2002; pp. 1423–1470. [\[CrossRef\]](#)
23. Daude, C.; Gutiérrez, H.; Melguizo, Á. What Drives Tax Morale? 2012. Available online: https://www.oecd-ilibrary.org/development/what-drives-tax-morale_5k8zk8m61kzq-en (accessed on 2 June 2022).
24. Mittone, L.; Patelli, P. Imitative Behaviour in Tax Evasion. In *Advances in Computational Economics*; Springer: New York, NY, USA, 2000; pp. 133–158. [\[CrossRef\]](#)
25. Davis, J.S.; Hecht, G.; Perkins, J.D. Social Behaviors, Enforcement, and Tax Compliance Dynamics. *Account. Rev.* **2003**, *78*, 39–69. [\[CrossRef\]](#)
26. Hokamp, S. Dynamics of tax evasion with back auditing, social norm updating, and public goods provision—An agent-based simulation. *J. Econ. Psychol.* **2014**, *40*, 187–199. [\[CrossRef\]](#)
27. Charteris, P.; Golden, B.; Garrick, D.J. Livestock breeding industries as complex adaptive systems. In Proceedings of the Conference of the Association for the Advancement of Animal Breeding and Genetics, Townsville, Australia, 26–28 July 2001; Volume 14, pp. 461–464.
28. Breckling, B. Individual-Based Modelling Potentials and Limitations. *Sci. World J.* **2002**, *2*, 1044–1062. [\[CrossRef\]](#)
29. Ligmann, A.; Sun, L. Applying time-dependent variance-based global sensitivity analysis to represent the dynamics of an agent-based model of land use change. *Int. J. Geogr. Inf. Sci.* **2010**, *24*, 1829–1850. [\[CrossRef\]](#)
30. Malleson, N.; Heppenstall, A.; See, L. Crime reduction through simulation: An agent-based model of burglary. *Comput. Environ. Urban Syst.* **2010**, *34*, 236–250. [\[CrossRef\]](#)
31. Lengnick, M. Agent-based macroeconomics: A baseline model. *J. Econ. Behav. Organ.* **2013**, *86*, 102–120. [\[CrossRef\]](#)
32. Conte, R.; Paolucci, M. On agent-based modeling and computational social science. *Front. Psychol.* **2014**, *5*, 668. [\[CrossRef\]](#)
33. Ghorbani, A.; Dechesne, F.; Dignum, V.; Jonker, C. Enhancing ABM into an Inevitable Tool for Policy Analysis. *J. Policy Complex Syst.* **2014**, *1*, 61–76. [\[CrossRef\]](#)
34. Siegfried, R. *Modeling and Simulation of Complex Systems*; Springer Fachmedien Wiesbaden: Wiesbaden, Germany, 2014. [\[CrossRef\]](#)
35. Tarvid, A. Complex Adaptive Systems and Agent-Based Modelling. In *Agent-Based Modelling of Social Networks in Labour-Education Market System*; Springer: Cham, Switzerland, 2016; pp. 23–38. [\[CrossRef\]](#)
36. Li, X.; Mao, W.; Zeng, D.; Wang, F.Y. Agent-Based Social Simulation and Modeling in Social Computing. In *Intelligence and Security Informatics*; Springer: Berlin/Heidelberg, Germany, 2008; pp. 401–412. [\[CrossRef\]](#)
37. Baldwin, W.C.; Sauser, B.; Cloutier, R. Simulation Approaches for System of Systems: Events-based versus Agent Based Modeling. *Procedia Comput. Sci.* **2015**, *44*, 363–372.
38. Lee, J.S.; Filatova, T.; Ligmann-Zielinska, A.; Hassani-Mahmoei, B.; Stonedahl, F.; Lorscheid, I.; Voinov, A.; Polhill, J.G.; Sun, Z.; Parker, D.C. The Complexities of Agent-Based Modeling Output Analysis. *J. Artif. Soc. Soc. Simul.* **2015**, *18*, 4. [\[CrossRef\]](#)
39. Axtell, R.L. Why Agents? On the Varied Motivations for Agent Computing in the Social Sciences. In *Workshop on Agent Simulation: Applications, Models, and Tools*, 2000. Available online: <http://www.brook.edu/dybdocroot/es/dynamics/papers/agents/agents.htm> (accessed on 2 June 2022).
40. Nourqolipour, R.; Shariff, R. How agent based modeling (ABM) can be linked to GIS for modelling land use and land cover change. In Proceedings of the MRSS 6th International Remote Sensing and GIS Conference and Exhibition, Kuala Lumpur, Malaysia, 28–29 April 2010.
41. Rand, W.; Rust, R.T. Agent-based modeling in marketing: Guidelines for rigor. *Int. J. Res. Mark.* **2011**, *28*, 181–193. [\[CrossRef\]](#)

42. Miller, M.Z.; Griendling, K.; Mavris, D.N. Exploring human factors effects in the Smart Grid system of systems Demand Response. In Proceedings of the 2012 7th International Conference on System of Systems Engineering (SoSE), Genova, Italy, 16–19 July 2012; pp. 1–6. [\[CrossRef\]](#)
43. Kostadinov, F.; Holm, S.; Steubing, B.; Thees, O.; Lemm, R. Simulation of a Swiss wood fuel and roundwood market: An explorative study in agent-based modeling. *For. Policy Econ.* **2014**, *38*, 105–118. [\[CrossRef\]](#)
44. Wilensky, U.; Rand, W. *An Introduction to Agent-Based Modeling: Modeling Natural, Social, and Engineered Complex Systems with NetLogo*; MIT Press: Cambridge, MA, USA, 2015.
45. Broniec, W.; An, S.; Rugarber, S.; Goel, A.K. Guiding Parameter Estimation of Agent-Based Modeling Through Knowledge-based Function Approximation. In Proceedings of the AAAI 2021 spring symposium on combining machine learning and knowledge engineering (AAAI-MAKE 2021), Palo Alto, CA, USA, 22–24 March 2021.
46. Tesfatsion, L. Agent-Based Computational Economics: Growing Economies From the Bottom Up. *Artif. Life* **2002**, *8*, 55–82.
47. Wang, C.; Hu, M.; Yang, L.; Zhao, Z. Prediction of air traffic delays: An agent-based model introducing refined parameter estimation methods. *PLoS ONE* **2021**, *16*, e0249754. [\[CrossRef\]](#)
48. Saeedi, S. Integrating macro and micro scale approaches in the agent-based modeling of residential dynamics. *Int. J. Appl. Earth Obs. Geoinf.* **2018**, *68*, 214–229. [\[CrossRef\]](#)
49. Sánchez-Marroño, N.; Alonso-Betanzos, A.; Fontenla-Romero, O.; Polhill, J.; Craig, T. Empirically-derived behavioral rules in agent-based models using decision trees learned from questionnaire data. In *Understanding Complex Systems*; Springer: Cham, Switzerland, 2017; pp. 53–76. [\[CrossRef\]](#)
50. Pouladi, P.; Afshar, A.; Molajou, A.; Afshar, M. Socio-hydrological framework for investigating farmers' activities affecting the shrinkage of Urmia Lake; hybrid data mining and agent-based modelling. *Hydrol. Sci. J.* **2020**, *65*, 1249–1261. [\[CrossRef\]](#)
51. Jäger, G. Replacing rules by neural networks a framework for agent-based modelling. *Big Data Cogn. Comput.* **2019**, *3*, 51. [\[CrossRef\]](#)
52. Romero-Mujalli, D.; Cappelletto, J.; Herrera, E.; Tárano, Z. The effect of social learning in a small population facing environmental change: An agent-based simulation. *J. Ethol.* **2017**, *35*, 61–73. [\[CrossRef\]](#)
53. Wolf, I.; Schröder, T.; Neumann, J.; de Haan, G. Changing minds about electric cars: An empirically grounded agent-based modeling approach. *Technol. Forecast. Soc. Chang.* **2015**, *94*, 269–285. [\[CrossRef\]](#)
54. van der Hoog, S. Surrogate Modelling in (and of) Agent-Based Models: A Prospectus. *Comput. Econ.* **2019**, *53*, 1245–1263. [\[CrossRef\]](#)
55. Neri, F. Combining machine learning and agent based modeling for gold price prediction. *Commun. Comput. Inf. Sci.* **2019**, *900*, 91–100. [\[CrossRef\]](#)
56. Cockrell, C.; An, G. Utilizing the Heterogeneity of Clinical Data for Model Refinement and Rule Discovery Through the Application of Genetic Algorithms to Calibrate a High-Dimensional Agent-Based Model of Systemic Inflammation. *Front. Physiol.* **2021**, *12*, 726. [\[CrossRef\]](#)
57. Ye, P.; Chen, Y.; Zhu, F.; Lv, Y.; Lu, W.; Wang, F. Bridging the Micro and Macro: Calibration of Agent-Based Model Using Mean-Field Dynamics. *IEEE Trans. Cybern.* **2021**, *52*, 11397–11406. [\[CrossRef\]](#)
58. Kim, D.; Yun, T.S.; Moon, I.C.; Bae, J. Automatic calibration of dynamic and heterogeneous parameters in agent-based models. *Auton. Agents Multi-Agent Syst.* **2021**, *35*, 1–66. [\[CrossRef\]](#)
59. Rajabi, M.; Pilesjö, P.; Shirzadi, M.; Fadaei, R.; Mansourian, A. A spatially explicit agent-based modeling approach for the spread of Cutaneous Leishmaniasis disease in central Iran, Isfahan. *Environ. Model. Softw.* **2016**, *82*, 330–346. [\[CrossRef\]](#)
60. Hayashi, S.; Prasasti, N.; Kanamori, K.; Ohwada, H. Improving behavior prediction accuracy by using machine learning for agent-based simulation. *Lect. Notes Comput. Sci. (Incl. Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinform.)* **2016**, *9621*, 280–289. [\[CrossRef\]](#)
61. Vahdati, A.R.; Weissmann, J.; Timmermann, A.; Ponce de León, M.; Zollikofer, C. Drivers of Late Pleistocene human survival and dispersal: An agent-based modeling and machine learning approach. *Quat. Sci. Rev.* **2019**, *221*, 105867. [\[CrossRef\]](#)
62. Ozik, J.; Collier, N.; Wozniak, J.; Macal, C.; An, G. Extreme-scale dynamic exploration of a distributed agent-based model with the EMEWS framework. *IEEE Trans. Comput. Soc. Syst.* **2018**, *5*, 884–895. [\[CrossRef\]](#)
63. Chen, S.; Londoño-Larrea, P.; McGough, A.; Bible, A.; Gunaratne, C.; Araujo-Granda, P.; Morrell-Falvey, J.; Bhowmik, D.; Fuentes-Cabrera, M. Application of Machine Learning Techniques to an Agent-Based Model of Pantoea. *Front. Microbiol.* **2021**, *12*, 2638. [\[CrossRef\]](#)
64. Perry, G.; O'Sullivan, D. Identifying Narrative Descriptions in Agent-Based Models Representing Past Human-Environment Interactions. *J. Archaeol. Method Theory* **2018**, *25*, 795–817. [\[CrossRef\]](#)
65. Garg, A.; Yuen, S.; Seekhao, N.; Yu, G.; Karwowski, J.; Powell, M.; Sakata, J.; Mongeau, L.; Jaja, J.; Li-Jessen, N. Towards a physiological scale of vocal fold agent-based models of surgical injury and repair: Sensitivity analysis, calibration and verification. *Appl. Sci.* **2019**, *9*, 2974. [\[CrossRef\]](#)
66. Edali, M.; Yücel, G. Exploring the behavior space of agent-based simulation models using random forest metamodels and sequential sampling. *Simul. Model. Pract. Theory* **2019**, *92*, 62–81. [\[CrossRef\]](#)
67. Gursoy, F.; Badur, B. An Agent-Based Modeling Approach to Brain Drain. *IEEE Trans. Comput. Soc. Syst.* **2021**, *9*, 356–365. [\[CrossRef\]](#)

68. Xu, T.; Gao, J.; Coco, G.; Wang, S. Urban expansion in Auckland, New Zealand: A GIS simulation via an intelligent self-adapting multiscale agent-based model. *Int. J. Geogr. Inf. Sci.* **2020**, *34*, 2136–2159. [\[CrossRef\]](#)
69. Xiao, S.; Liu, R. Studies of covid-19 outbreak control using agent-based modeling. *Complex Syst.* **2021**, *30*, 297–321. [\[CrossRef\]](#)
70. Zhang, Y.; Grignard, A.; Lyons, K.; Aubuchon, A.; Larson, K. Real-time machine learning prediction of an agent-based model for urban decision-making (extended abstract). In Proceedings of the AAMAS International Foundation for Autonomous Agents and Multiagent Systems (IFAAMAS), Stockholm, Sweden, 10–15 July 2018; Volume 3, pp. 2171–2173.
71. Ten Broeke, G.; van Voorn, G.; Ligtenberg, A.; Molenaar, J. The Use of Surrogate Models to Analyse Agent-Based Models. *J. Artif. Soc. Soc. Simul.* **2021**, *24*, 3. [\[CrossRef\]](#)
72. Janssen, S.; Sharpanskykh, A.; Curran, R.; Langendoen, K. Using causal discovery to analyze emergence in agent-based models. *Simul. Model. Pract. Theory* **2019**, *96*, 101940. [\[CrossRef\]](#)
73. Nunes, A.; Zwick, M.; Wakeland, W. Sensitivity analysis of an agent-based simulation model using reconstructability analysis. *Int. J. Gen. Syst.* **2021**, *50*, 319–338. [\[CrossRef\]](#)
74. Koda, H.; Arai, Z.; Matsuda, I. Agent-based simulation for reconstructing social structure by observing collective movements with special reference to single-file movement. *PLoS ONE* **2020**, *15*, e0243173. [\[CrossRef\]](#)
75. Esmaeili Aliabadi, D.; Kaya, M.; Sahin, G. Competition, risk and learning in electricity markets: An agent-based simulation study. *Appl. Energy* **2017**, *195*, 1000–1011. [\[CrossRef\]](#)
76. Dehghanpour, K.; Hashem Nehrir, M.; Sheppard, J.; Kelly, N. Agent-Based Modeling of Retail Electrical Energy Markets with Demand Response. *IEEE Trans. Smart Grid* **2018**, *9*, 3465–3475. [\[CrossRef\]](#)
77. Aghaie, A.; Heidary, M. Simulation-based optimization of a stochastic supply chain considering supplier disruption: Agent-based modeling and reinforcement learning. *Sci. Iran.* **2019**, *26*, 3780–3795. [\[CrossRef\]](#)
78. Ibrahim, M.; Hashmi, U.; Nabeel, M.; Imran, A.; Ekin, S. Embracing Complexity: Agent-based Modeling for HetNets Design and Optimization via Concurrent Reinforcement Learning Algorithms. *IEEE Trans. Netw. Serv. Manag.* **2021**, *18*, 4042–4062. [\[CrossRef\]](#)
79. Schauder, S.; Thomsen, M.; Nayga, R., Jr. Agent-based modeling insights into the optimal distribution of the Fresh Fruit and Vegetable Program. *Prev. Med. Rep.* **2020**, *20*, 101173. [\[CrossRef\]](#)
80. Harati, S.; Perez, L.; Molowny-Horas, R. Promoting the emergence of behavior norms in a principal-agent problem—An agent-based modeling approach using reinforcement learning. *Appl. Sci.* **2021**, *11*, 8368. [\[CrossRef\]](#)
81. Liang, Y.; Guo, C.; Ding, Z.; Hua, H. Agent-Based Modeling in Electricity Market Using Deep Deterministic Policy Gradient Algorithm. *IEEE Trans. Power Syst.* **2020**, *35*, 4180–4192. [\[CrossRef\]](#)
82. Sert, E.; Bar-Yam, Y.; Morales, A. Segregation dynamics with reinforcement learning and agent based modeling. *Sci. Rep.* **2020**, *10*, 11771. [\[CrossRef\]](#)
83. Jäger, G. Using Neural Networks for a Universal Framework for Agent-based Models. *Math. Comput. Model. Dyn. Syst.* **2021**, *27*, 162–178. [\[CrossRef\]](#)
84. Salle, I. Modeling expectations in agent-based models - An application to central bank's communication and monetary policy. *Econ. Model.* **2015**, *46*, 130–141. [\[CrossRef\]](#)
85. Dehghanpour, K.; Nehrir, M.; Sheppard, J.; Kelly, N. Agent-Based Modeling in Electrical Energy Markets Using Dynamic Bayesian Networks. *IEEE Trans. Power Syst.* **2016**, *31*, 4744–4754. [\[CrossRef\]](#)
86. Norman, M.; Koehler, M.; Kutarnia, J.; Silvey, P.; Tolk, A.; Tracy, B. Applying Complexity Science with Machine Learning, Agent-Based Models, and Game Engines: Towards Embodied Complex Systems Engineering. In *Proceedings of the Unifying Themes in Complex Systems IX*; Springer: Berlin/Heidelberg, Germany, 2018; pp. 173–183. [\[CrossRef\]](#)
87. Fuller, D.; de Arruda, E.; Ferreira Filho, V. Learning-agent-based simulation for queue network systems. *J. Oper. Res. Soc.* **2020**, *71*, 1723–1739. [\[CrossRef\]](#)
88. Cummings, P.; Crooks, A. Development of a Hybrid Machine Learning Agent Based Model for Optimization and Interpretability. In *International Conference on Social Computing, Behavioral-Cultural Modeling and Prediction and Behavior Representation in Modeling and Simulation*; Springer: Cham, Switzerland, 2020; pp. 151–160. [\[CrossRef\]](#)
89. Grimm, V.; Berger, U.; Bastiansen, F.; Eliassen, S.; Ginot, V.; Giske, J.; Goss-Custard, J.; Grand, T.; Heinz, S.K.; Huse, G.; et al. A standard protocol for describing individual-based and agent-based models. *Ecol. Model.* **2006**, *198*, 115–126. [\[CrossRef\]](#)
90. Chakraborti, A.; Patriarca, M. Gamma-distribution and wealth inequality. *Pramana* **2008**, *71*, 233–243. [\[CrossRef\]](#)
91. Pinder, J.E.; Wiener, J.G.; Smith, M.H. The Weibull Distribution: A New Method of Summarizing Survivorship Data. *Ecology* **1978**, *59*, 175–179. [\[CrossRef\]](#)
92. CONAPO. Datos Abiertos. Indicadores Demográficos 1950–2050. 2018. Available online: http://www.conapo.gob.mx/work/models/CONAPO/Datos_Abiertos/Proyecciones2018/ind_dem_proyecciones.csv (accessed on 21 October 2021).
93. INEGI. Encuesta Nacional de Ocupación y Empleo (ENOE). 2019. Available online: <https://www.inegi.org.mx/programas/enoe/15ymas/> (accessed on 27 September 2020).
94. Bonet, J.A.; Rueda, F. *Esfuerzo Fiscal en los Estados Mexicanos*; IDB Publications (Working Papers) 3946; Inter-American Development Bank: Washington, DC, USA, 2012.
95. Witten, I.; Frank, E.; Hall, M.; Pal, C. *Data Mining: Practical Machine Learning Tools and Techniques*; The Morgan Kaufmann Series in Data Management Systems; Elsevier Science: Amsterdam, The Netherlands, 2017.
96. Lisic, J.; Cruze, N. Local Pivotal Methods for Large Surveys. In Proceedings of the International Conference on Establishment Surveys, Geneva, Switzerland, 20–23 June 2016.

97. Gościk, J.; Łukaszuk, T. Application of the recursive feature elimination and the relaxed linear separability feature selection algorithms to gene expression data analysis. *Adv. Comput. Sci. Res.* **2013**, *10*, 39–52.
98. Kuhn, M. Building Predictive Models in R Using the caret Package. *J. Stat. Softw.* **2008**, *28*, 1–26. [[CrossRef](#)]
99. Wright, M.N.; Ziegler, A. ranger: A Fast Implementation of Random Forests for High Dimensional Data in C++ and R. *J. Stat. Softw.* **2017**, *77*, 1–17. [[CrossRef](#)]
100. Han, S.; Kim, H. Optimal Feature Set Size in Random Forest Regression. *Appl. Sci.* **2021**, *11*, 3428. [[CrossRef](#)]
101. INEGI. Medición de la Economía Informal. 2019. Available online: <https://www.inegi.org.mx/temas/pibmed/> (accessed on 23 October 2021).
102. Samaniego-Breach, N. La participación del trabajo en el ingreso nacional: El regreso a un tema olvidado. *Econ. UNAM* **2014**, *11*, 52–77. [[CrossRef](#)]
103. Wilensky, U. NetLogo. Center for Connected Learning and Computer-Based Modeling, 1999. Available online: <http://ccl.northwestern.edu/netlogo/> (accessed on 12 January 2022).
104. Marks, R.E. Validation and model selection: Three similarity measures compared. *Complex. Econ.* **2013**, *2*, 41–61. [[CrossRef](#)]
105. Fagiolo, G.; Guerini, M.; Lamperti, F.; Moneta, A.; Roventini, A. Validation of Agent-Based Models in Economics and Finance. In *Simulation Foundations, Methods and Applications*; Springer International Publishing: Berlin/Heidelberg, Germany, 2019; pp. 763–787. [[CrossRef](#)]
106. Windrum, P.; Fagiolo, G.; Moneta, A. Empirical Validation of Agent-Based Models: Alternatives and Prospects. *J. Artif. Soc. Soc. Simul.* **2007**, *10*, 8.
107. Joanes, D.N.; Gill, C.A. Comparing Measures of Sample Skewness and Kurtosis. *J. R. Stat. Society. Ser. D (Stat.)* **1998**, *47*, 183–189. [[CrossRef](#)]
108. INEGI. Finanzas Públicas Estatales y Municipales. 2019. Available online: <https://www.inegi.org.mx/programas/finanzas/> (accessed on 13 September 2021).
109. Akoglu, H. User's guide to correlation coefficients. *Turk. J. Emerg. Med.* **2018**, *18*, 91–93. [[CrossRef](#)]
110. Xu, S.; Chen, M.; Feng, T.; Zhan, L.; Zhou, L.; Yu, G. Use ggbreak to Effectively Utilize Plotting Space to Deal With Large Datasets and Outliers. *Front. Genet.* **2021**, *12*, 2122. [[CrossRef](#)]

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.



Contents lists available at ScienceDirect

Applied Mathematical Modelling

journal homepage: www.elsevier.com/locate/apm

Discriminative learning of bayesian network parameters by differential evolution

Alejandro Platas-López*, Efrén Mezura-Montes, Nicandro Cruz-Ramírez, Alejandro Guerra-Hernández

Instituto de Investigaciones en Inteligencia Artificial, Universidad Veracruzana, Sebastián Camacho No. 5, Col. Centro, C.P. 91000, Xalapa, Veracruz, México

ARTICLE INFO

Article history:

Received 2 April 2020

Revised 28 October 2020

Accepted 11 December 2020

Available online 23 December 2020

Keywords:

Bayesian networks

Differential evolution

Discriminative learning

ABSTRACT

This work proposes Differential Evolution (DE) to train parameters of Bayesian Networks (BN) for optimizing the Conditional Log-Likelihood (Discriminative Learning) instead of the log-likelihood (Generative Learning). Any given BN structure encodes assumptions about conditional independencies among the attributes and will result in error if they do not hold in the data. Such an error includes the classification dimension. The main goal of Discriminative learning is minimize this type of error. In this sense, although BN Discriminative Parameter Learning algorithms have been proposed, to the best of the authors' knowledge, a meta-heuristic approach has not been devised yet. Thus, this is our main contribution: to come up with this kind of solution and evaluate its behavior so that its feasibility in this domain can be determined. Regarding the proposed method, the bias provided by the best solution in the population improves DE's performance. DE variants based on JADE, such as L-SHADE and C-JADE, are especially recommended when introducing adaptation mechanisms of mutation and crossover parameters thus reducing the dependence on their calibration. L-SHADE is computationally recommended over any other DE variant. DE approach works well in essentially every standard situation, so DE approach is robust and at least as good, and often better, than the state-of-art method for Discriminative Learning called WANBIA. Our results suggest a potential benefit for Discriminative parameter learning with strong independence assumptions among attributes and that it typically produces more accurate classifiers than generative learning.

© 2020 Elsevier Inc. All rights reserved.

1. Introduction

According to Friedman [1], a Bayesian Network (BN) is a graphical model that encodes a joint probability distribution over a set of discrete random variables. A BN B for a set of discrete random variables $\mathbf{X} = \{X_1, \dots, X_n\}$ is defined as a pair of components $B = \langle G, \Theta \rangle$, where G denotes a directed acyclic graph whose vertices correspond to these variables and whose edges represent direct dependencies among variables. The Markov condition establishes that a variable is independent of its non-descendants given its parents in the graph. Θ denotes a set of parameters $\theta_{x_i | \pi_{x_i}}$ for each possible value x_i of X_i and

* Corresponding author.

E-mail addresses: zs19019827@estudiantes.uv.mx (A. Platas-López), emezura@uv.mx (E. Mezura-Montes), ncruz@uv.mx (N. Cruz-Ramírez), aguerra@uv.mx (A. Guerra-Hernández).<https://doi.org/10.1016/j.apm.2020.12.026>

0307-904X/© 2020 Elsevier Inc. All rights reserved.

Π_{x_i} corresponds to the set of parents of X_i in G . In this way, a BN defines a unique probability distribution over $\mathbf{X}$ given by

$$P_B(X_1, \dots, X_n) = \prod_{i=1}^n \theta_{x_i} | \Pi_{x_i}$$

Evidently, the difficulty of building a BN resides in deciding the specific graph structure and the corresponding parameter values for a given problem. Traditionally this is done by exploiting domain expertise, but it is possible to adopt a data-driven inductive approach: Both components, structure and parameters, can be learned from data, i.e., given a training set $S = \{\mathbf{x}_1, \dots, \mathbf{x}_N\}$ of instances of $\mathbf{X}$, find a network B that *best matches* S [1]. This work adopts the data-driven approach to learn the parameters of BN B .

Two paradigms are distinguished for this: Generative Learning optimizes Log-Likelihood (LL) in order to obtain the parameters that characterize the joint distribution in the form of local conditional distributions, and subsequently estimates class conditional probabilities using the Bayes rule. Although this paradigm is computationally efficient, it is likely to generate overfitting classifiers [2]. On the other hand, Discriminative Learning optimizes Conditional Log-Likelihood (CLL) to directly estimate the parameters associated with conditional class distribution, generating classifiers that tend to minimize the classification error. In addition, the effect caused by the assumption of conditional independence among attributes in the network structure, but which may be violated in the data, is reduced. However, the huge search space defined by the parameters that optimize the CLL function motivates this work to find efficient and effective search algorithms for Discriminative Learning of parameters in BN classifiers [1,3].

Different algorithms have been developed with the purpose of generating unbiased classifiers that mitigate the assumption of conditional independence among attributes. In this paper, we propose the use of Differential Evolution (DE) algorithms [4], a kind of evolutionary algorithms, for learning the parameters of BN, optimizing CLL. Among the evolutionary approaches, DE algorithms, has emerged as one of the most frequently used algorithms for solving complex optimization problems due to its efficiency, flexibility and versatility [5,6]. To the best of our knowledge, the use of DE for Discriminative Learning of Bayesian Network parameters is currently unexplored. So, the aim of this work is to understand the behavior of DE in this particular optimization task, using different structures learned in turn by Naïve Bayes [7], Tree-Augmented Network [1] and hill climbing greedy search [8].

The rest of the paper is organized as follows. Section 2 introduces background and related work on learning the BN structures and parameters. Three methods for learning structures adopted in this work are briefly described: Naïve Bayes, Tree Augmented Naïve Bayes, and Hill Climbing Greedy Search. Then, different works on the Generative and Descriptive Learning of BN parameters are discussed. Section 3 introduces Differential Evolution and its variants considered in this work. It also defines the general framework for learning the parameters of BN following this approach. Section 4 defines the experimental setup, datasets where experiments were conducted are presented, calibration of parameters for DE variants is also presented and the tests performed to assess the performance of each algorithm is described. Results are presented in Section 5, performance in terms of the quality of the CLL solutions between the proposed approach and the state-of-the-art heuristic approach on Discriminative Learning of parameters is compared. A comparison in terms of predictive accuracy is included, where generative learning algorithms are also compared. Section 6 concludes this paper and gives possible paths of future work.

2. Background

Regarding the structure, probably the most common BN model is naïve Bayes [9]. In this model, the class variable, i.e., the one to be predicted, is the root of the structure and the attribute variables are leaves. The model is said to be *naïve* because it assumes that the attributes are conditionally independent of each other, given the class. This independence simplifies the representation of knowledge by reducing the number parameters Θ ; and reasoning, by the propagation of probabilities. All this results in a model that, regardless the strong assumptions of independence between attributes, turns out to do surprisingly well in a wide range of applications and scale very well to very large problems. However, biased classifiers are obtained if such independence is not present in data [2].

There are different approaches to achieve efficient heuristics for learning the structure of BN [10]. Some of them guide the search by mitigating the naïve assumption of independence between the attributes, such as the case of the Tree Augmented Naïve Bayes (TAN) algorithm [1], which is a semi-naïve Bayesian learning method. It employs a tree structure, in which each attribute only depends on the class and another attribute. To perform the classification, a maximum weighted spanning tree is used that maximizes the underlying probability distribution of the training data [11]. TAN uses a metric to score BN for searching, e.g., the minimal description length (MDL). Searching starts with the empty network and successively applies local operations that maximally improve the score until a local minima is found. The operations applied by the search procedure include arc addition, arc deletion, and arc reversal. As a classifier, TAN improves naïve Bayes in terms of bias at the cost of an increase in variance due to the growth in the number of parameters to estimate. Although the BN learned with TAN usually outperforms naïve Bayes, the former performs worse when coping with data sets of many attributes. Friedman [1] explains this behavior given the definition of MDL and other similar metrics, i.e., measuring the error of the learned BN over all the variables in the domain, including the class variable. Minimizing this error, however, does not necessarily minimize the local error in predicting the class variable given the attributes, especially in cases where there are many attributes.

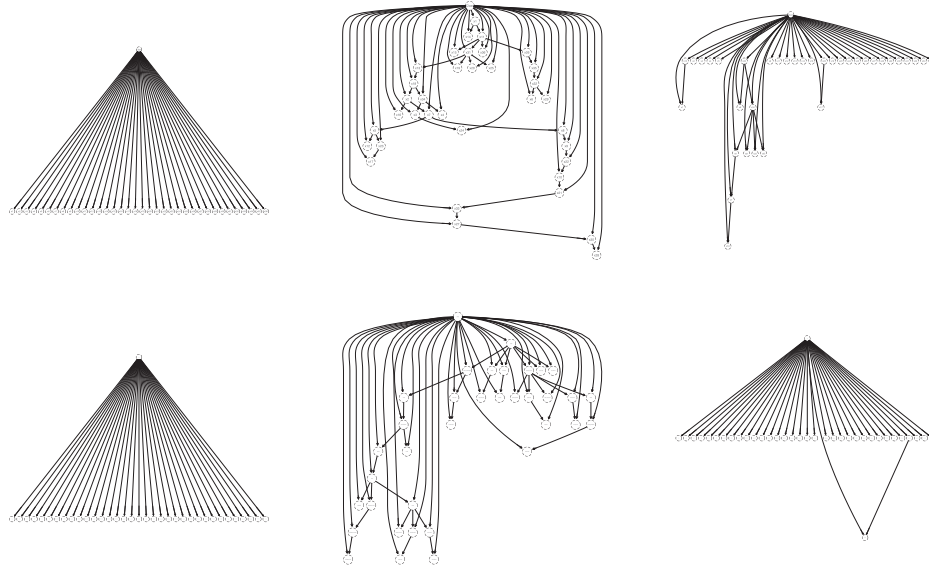


Fig. 1. Bayesian network structures. Trained, from left to right column, with Naïve Bayes, TAN and HCS on data sets Chess (first row) and Soybean large (second row).

Keogh and Pazzani [8] take a different approach to construct augmented BN. Rather than directly attempting to approximate the underlying probability distribution, they concentrate solely on using the same representation to improve classification accuracy with a hill climbing greedy search (HCS): The BN is initialized to naïve Bayes, initializing a set of orphans O , i.e., nodes without a parent, other than the class node, as the full set of attribute nodes $X_1, X_2, \dots, X_N$. Each possible arc from X_i to X_j ($X_i \neq X_j, X_j \in O$) is evaluated using leave-one-out cross validation to estimate the accuracy of the network with that arc added. If no arc produces an improvement in classification accuracy, the current classifier is returned; Otherwise, the arc that gives the highest improvement is retained and the node pointed by the arc is removed from O . This process is repeated until O contains just one node, or there are no arcs that can be added to improve classification accuracy. Interestingly, these authors demonstrate that the best approximation of the underlying probability distribution does not necessarily result in the most accurate classifier.

To intuitively show how those three different algorithms to learn structures work, Fig. 1 is plotted. While naïve Bayes does strong assumptions about independence among attributes, it produces networks no connected among attributes. TAN tends to generate highly connected networks. Networks produced by HCS are few connected. These three types of networks will help you understand the challenges optimization algorithms face as the search space becomes complex.

Regarding the parameters, Generative Learning of Θ is based on two steps, the first involves the maximization of $P(Y, \mathbf{X})$, where Y is the random variable associated with class label and $\mathbf{X}$ is the set of random variables associated with attributes; and the second step is the application of the Bayes rule to obtain $P(Y|\mathbf{X})$. Sundararajan and Mengshoel [12] proposed optimizing LL (generative approach) with a Genetic Algorithm combined with Expectation Maximization (GAEM). Their proposal combines the global search and local search properties of the respective algorithms. Part of notation and definitions used throughout this paper are taken from that work.

In Discriminative Learning, it is possible to directly optimize $P(Y|\mathbf{X})$, maximizing the CLL. For classification purposes, this uses to be a more effective objective function [1,2]. Zaidi and colleagues [2] define CLL as:

$$CLL(B) = \sum_{j=1}^N \log P_B(Y^{(j)} | \mathbf{X}^{(j)}), \quad (1)$$

which is re-defined as:

$$= \left(\log \theta_Y + \sum_{i=1}^n \log \theta_{X_i | Y, \Pi_i(\mathbf{X})} \right) - \log \sum_{Y'} \left(\theta_{Y'} \prod_{i=1}^n \theta_{X_i | Y', \Pi_i(\mathbf{X})} \right). \quad (2)$$

The objective function is to maximize Eq. 2 optimizing parameters (θ). Two explanations on why maximizing CLL can provide a better model of conditional distribution are given in [2]. First, parameters are allowed to be set in such a way that the assumption of conditional independence between attributes that is represented in the structure G but may not be present in the data set S is reduced. Second, given that $LL(B) = CLL(B) + LL(B)$, optimizing $LL(B)$ will lead to poor estimates for classification if $CLL(B) \ll LL(B)$, i.e., when $LL(B)$ dominates.

Discriminative Learning of BN parameters is not a new idea. Some researchers have made contributions, but none from a meta-heuristic approach. Greiner et al. [13–15] presented a simple gradient-descent algorithm, Extended Logistic Regression (ELR), a discriminative parameter-learning algorithm that maximizes conditional likelihood (CL) for a fixed BN structure. ELR extends Logistic Regression by computing the parameters that maximize CL. They demonstrated that it often produces classifiers that are more accurate than the ones produced using the generative approach which finds maximal likelihood parameters. That performance was especially true when learning parameters for incorrect structures, such as Naïve Bayes. ELR attempts to optimize the CLL of the belief net B given a sample S .

Friedman [1] note that maximizing this score will typically produce a classifier that comes close to minimizing the classification error. Unfortunately, the complexity of finding the Bayesian network parameters that optimize Eq. 1 is NP-hard [16]. According to [15], ELR is not guaranteed to find the parameters that optimize conditional likelihood.

Su et al. [3] proposed a discriminative parameter learning method, called Discriminative Frequency Estimate (DFE). In summary, DFE learns parameters by considering the likelihood information and the prediction error, and thus is actually a combination of generative and discriminative learning. However, the generative step dominates the discriminative one since if we eliminate the generative step, the performance of DFE is poor, while if we eliminate the discriminative step DFE at least performs like a traditional naïve Bayes. DFE and ELR have similar performance in term of accuracy but DFE is faster than ELR.

Pernkopf and Wohlmayr [17] present three discriminative Expectation-Maximization like parameter learning methods for Bayesian network classifiers called extended Baum-Welch (EBW) for Bayesian Networks, exact CL decomposition (ECL) and approximate CL decomposition (ACL). ECL and ACL iteratively optimize the parameters (initialized to ML) with a 2-step algorithm. In the first step, either the class posterior probabilities or class assignments are determined based on current parameter estimates. Based on these posteriors (class assignment, respectively), the parameters are updated in the second step. Although ECL and ACL are computationally efficient, their performance is worse than a conjugate gradient conditional likelihood (CGCL) parameter optimization algorithm as in [15] when network structures are not optimized for classification. EBW is relatively competitive compared to [15].

Zaidi and colleagues [18] proposed a simple heuristic approach, called Weighting attributes to Alleviate Naïve Bayes' Independence Assumption (WANBIA) and extend it [2] to any BN classifier. This is a two-step learning process: First, a Generative Learning step in which LL is the objective function for all local conditional distributions in B ; Second, a Discriminative Learning step, where each parameter learned in the generative step is associated with a weight to re-parameterize the class-conditional distribution in terms of those weights when optimizing CLL. The benefit of their technique progressively increases when features not contribute uniformly to the class prediction, which is the case for the vast majority of real-world datasets. To our knowledge, WANBIA is the state-of-the-art algorithm for Discriminative Learning of parameters for Bayesian network classifiers.

Observe that if the BN structure encodes correctly the independence among attributes, both approaches should lead to the same results. So, structures in which there is a strong assumption of independence, like in Naïve Bayes, are expected to benefit more from the discriminative approach than in those where this assumption has been relaxed.

3. Differential evolution

Differential Evolution (DE) [4] has been extensively used to solve optimization problems in different domains [19]. It shares similarities with traditional Evolutionary Algorithms, as two main processes derive evolution: A perturbation process (crossover and mutation) guarantees the exploration of the search space, while a selection process guarantees the exploitation properties of DE [20].

DE algorithm [4] is a population based algorithm. Its population P is composed of NP vectors:

$$\mathbf{P}_g = (\bar{z}_{1,g}, \dots, \bar{z}_{i,g}, \dots, \bar{z}_{NP,g}), i = 1, 2, \dots, NP,$$

where g denotes a generation index, $g = 1, 2, \dots, G_{MAX}$. Each individual consists of D variables:

$$\bar{z}_{i,g} = (z_{i,1,g}, z_{i,2,g}, \dots, z_{i,D,g}).$$

A randomly initialized population of NP vectors is evolved throughout G_{MAX} generations guiding the vectors in a search space toward promising regions. At the end of the evolutionary process, i.e., after G_{MAX} generations, DE stops and returns the best fitted vector as the final solution. During each generation, DE employs three operations for each individual, namely mutation, crossover and selection.

Mutation: A mutant vector $\bar{v}_{i,g}$ is created using one of the mutation strategies. Depending on the strategy, different vectors are selected and their difference multiplied by scale factor F is added to other vector. Each strategy has a different ability to maintain the population diversity which might increase/decrease the convergence rate during the evolutionary process. Subsection 3.1 presents different mutation strategies frequently used in the literature.

Table 1
Commonly used mutation strategies.

Nomenclature	Mutation strategy
rand/1	$\bar{v}_{i,g} = \bar{z}_{r_0,g} + F(\bar{z}_{r_1,g} - \bar{z}_{r_2,g})$
best/1	$\bar{v}_{i,g} = \bar{z}_{best,g} + F(\bar{z}_{r_1,g} - \bar{z}_{r_2,g})$
rand/2	$\bar{v}_{i,g} = \bar{z}_{best,g} + F(\bar{z}_{r_1,g} - \bar{z}_{r_2,g} + \bar{z}_{r_3,g} - \bar{z}_{r_4,g})$
best/2	$\bar{v}_{i,g} = \bar{z}_{best,g} + F(\bar{z}_{r_1,g} - \bar{z}_{r_2,g} + \bar{z}_{r_3,g} - \bar{z}_{r_4,g})$
current-to-rand/1	$\bar{v}_{i,g} = \bar{z}_{i,g} + K(\bar{z}_{r_2,g} - \bar{z}_{i,g}) + F(\bar{z}_{r_0,g} - \bar{z}_{r_1,g})$
current-to-best/1	$\bar{v}_{i,g} = \bar{z}_{i,g} + K(\bar{z}_{best,g} - \bar{z}_{i,g}) + F(\bar{z}_{r_0,g} - \bar{z}_{r_1,g})$
rand-to-best/1	$\bar{v}_{i,g} = \bar{z}_{i,g} + K(\bar{z}_{best,g} - \bar{z}_{r_0,g}) + F(\bar{z}_{r_1,g} - \bar{z}_{r_2,g})$

Crossover: Each vector in the current population (called target or parent at this time) and the mutant vector $\bar{v}_{i,g}$ generated by one of the mutation strategies are used in the next operation, called crossover, to create a trial vector $\bar{u}_{i,g}$. At this step, the Crossover Rate $CR \in [0, 1]$ parameter is used to control the influence of the mutant vector with respect to the parent (target) vector in the construction of the offspring vector.

Selection: After crossover, the trial vector is evaluated in the objective function $f(\bar{u}_{i,g})$, and then the selection operation compares two vectors, the parent (target) vector $\bar{z}_{i,g}$ and its corresponding trial vector $\bar{u}_{i,g}$, according to their objective function values. The better vector will become a member of the population for the next generation. The selection operation for a maximization optimization problem is defined as follow:

$$\bar{z}_{i,g+1} = \begin{cases} \bar{u}_{i,g}, & \text{if } f(\bar{u}_{i,g}) \geq f(\bar{z}_{i,g}) \\ \bar{z}_{i,g}, & \text{otherwise} \end{cases}$$

Summarizing, there are three parameters to be defined: The population size (NP), the crossover rate (CR) that controls the influence of a mutant vector in the contents of the offspring vector, and the scale factor F (or K depending the mutation strategy) that scales the vector differences to calculate the mutant values. DE executes the perturbation based on the distribution of the solutions of the current population. In this sense, possible search directions and the corresponding step size depend on the location of the selected individuals to compute the mutation values [20].

In what follows the different DE variants considered in this work are introduced, as well as the general setting for learning the parameters of a BN by DE.

3.1. DE Variants

Based on different strategies followed by the mutation operator, there are different DE variants, which define the way in which the mutant and trial vectors are generated. Beyond the proper control parameters setting, another important factor is the choice of a particular variant.

There is a nomenclature scheme developed to reference different DE variants. The most popular variant is called “DE/rand/1/bin”, where “DE” means Differential Evolution, the word “rand” indicates that the so-called base vector ($\bar{z}_{r_0,g}$) is randomly chosen, “1” is the number of vector pairs (i.e., vector differences to be calculated) chosen, and finally “bin” means that a binomial recombination is chosen. Table 1 summarise the most common mutation strategies.

Once the mutant vector $\bar{v}_{i,g}$ at generation g is generated, the crossover operation will be applied to generate the final trial/offspring vector $\bar{u}_{i,g}$. A common crossover operator is the Binomial one, which is defined as follows:

$$\bar{u}_{j,i,g} = \begin{cases} \bar{v}_{j,i,g} & \text{if } \text{rand}(0, 1) \leq CR_i \text{ or } j = j_{\text{rand}}, \\ \bar{z}_{j,i,g} & \text{otherwise} \end{cases}$$

where j represent the j^{th} component of the trial vector generated at random with a uniform distribution. Another common crossover operator is the Exponential crossover, defined as:

$$\bar{u}_{j,i,g} = \begin{cases} \bar{v}_{j,i,g} & \text{while } \text{rand}(0, 1) \leq CR_i \text{ or } j = j_{\text{rand}}, \\ \bar{z}_{j,i,g} & \text{otherwise} \end{cases}$$

JADE [19], a more recent DE variant, includes some mechanisms to decrease the dependence to its parameter values such as F and CR . JADE uses a mutation strategy called “DE/current-to-pbest”, where $p \in (0, 1]$. Vectors are selected from the best 100% vectors in the current population. The mutation procedure is described as:

$$\bar{v}_{i,g} = \bar{z}_{i,g} + F_i \cdot (\bar{z}_{best,g}^p - \bar{z}_{i,g}) + F_i \cdot (\bar{z}_{r_1,g} - \bar{z}_{r_2,g}) \quad (3)$$

where $\bar{z}_{best,g}^p$ is a randomly selected vector from the best 100% vectors. $\bar{z}_{r_1,g}$ is a randomly selected vector from the population. JADE, with the aim to maintaining diversity, uses an optional external archive to store the parent vectors that are replaced by the trial vectors. $\bar{z}_{r_2,g}$ is randomly selected from the union of the current population and the external archive. In the JADE version with Archive, the size of the external archive equals the size of the population (NP). When the size of the archive exceeds the threshold, a vector within it is randomly deleted to make space for the new one.

After generating the mutant vector $\bar{v}_{i,g}$, it is combined with the parent (target) $\bar{z}_{i,g}$ in order to generate the trial vector $\bar{u}_{i,g}$ following the binomial crossover operator. At each generation g , the crossover rate CR_i of each individual $\bar{z}_i$ is independently generated according to a Normal distribution with mean μ_{CR} and standard deviation 0.1

$$CR_i = \text{randn}_i(\mu_{CR}, 0.1)$$

and then truncated to $[0, 1]$. Denote S_{CR} as the set of all successful crossover probabilities CR_i at generation g . The mean μ_{CR} is initialised to be 0.5 and then updated at the end of each generation as

$$\mu_{CR} = (1 - c) \cdot \mu_{CR} + c \cdot \text{mean}_A(S_{CR})$$

where c is a positive constant between 0 and 1 and $\text{mean}_A(\cdot)$ is the usual arithmetic mean. Similarly, at each generation g , the scale factor F_i of each individual $\bar{z}_i$ is independently generated according to a Cauchy distribution with location parameter μ_F and scale parameter 0.1

$$F_i = \text{randc}_i(\mu_F, 0.1)$$

and then truncated to be 1 if $F_i \geq 1$ or regenerated if $F_i \leq 1$. Denote S_F as the set of all successful mutation factors in generation g . The location parameter μ_F of the Cauchy distribution is initialised to be 0.5 and then updated at the end of each generation as:

$$\mu_F = (1 - c) \cdot \mu_F + c \cdot \text{mean}_L(S_F)$$

where $\text{mean}_L(\cdot)$ is the Lehmer mean

$$\text{mean}_L(S_F) = \frac{\sum_{F \in S_F} F^2}{\sum_{F \in S_F} F}$$

Success-History based Adaptive DE (SHADE) is an improved JADE version. L-SHADE further extends SHADE with a Linear Population Size Reduction (LPSR), which decreases the population size according to a linear function [21]. The success-history based adaptation uses a historical memory M_{CR} , M_F which stores a set of CR and F values that have performed well in the past, and generates new CR , F pairs by directly sampling the parameter space close to one of those stored pairs.

L-SHADE maintains a historical memory with H entries for both DE control parameters, $CR \in [0, 1]$ and $F \in [0, 1]$, M_{CR} , M_F , respectively. At the beginning, the contents of $M_{CR,k}$, $M_{F,k}$ ($k = 1, \dots, H$) are all initialised to 0.5. In each generation g , the control parameters CR_i and F_i used by each individual $\bar{z}_i$ are generated by randomly selecting an index r_i from $[1, H]$, and then applying the formulas below:

$$CR_i = \begin{cases} 0 & \text{if } M_{CR,r_i} = \perp \\ \text{randn}_i(M_{CR,r_i}, 0.1) & \text{otherwise} \end{cases}$$

$$F_i = \text{randc}_i(M_{F,r_i}, 0.1)$$

In cases when F_i or CR_i generated values fall outside the range $[0, 1]$, the JADE procedure is applied. $\perp$ is a “terminal value”. The mutant vector $\bar{v}_{i,g}$ is generated from existing population members by applying the “DE/current-to-pbest/1” mutation strategy as in Eq. 3. To generate the trial vector $\bar{u}_{i,g}$, the binomial crossover operator is applied. L-SHADE adopts the external archive for maintaining diversity. Both, the mutation strategy and the external archive are borrowed from JADE.

Those CR_i and F_i values that generated a trial vector $\bar{u}_{i,g}$ which was better than its parent $\bar{z}_{i,g}$ are recorded as S_{CR} , S_F . At the end of the generation, the memory contents are updated. L-SHADE incorporates a Linear Population Size Reduction (LPSR) in order to improve the performance. After each generation g , the population size in the next generation, N_{g+1} , is computed, i.e., the worst ranked individuals are deleted from the population whenever $N_{g+1} < N_g$. The archive size $|A|$ is also adjusted to the corresponding N_{g+1} new population size.

Chaotic Local Search based on JADE (CJADE) [22] exploits the chaos' randomness and ergodicity properties to improve the search ability and alleviate the problem of falling into local optima. Its authors present four different CJADE variants. With a selection mechanism that combines 12 different chaotic maps based on memory, a variant called CJADE-M was proposed. That variant will be used in this research because it has a highly competitive performance reported in the literature [22].

CJADE-M uses $J = 12$ chaotic maps for Chaotic Local Search (CLS). In each generation g , CLS searches all dimensions of the current global best individual $\bar{z}_{best,g}$ in order to generate a new individual $\bar{z}_{best',g}$ as follows:

$$\bar{z}_{best',g} = \bar{z}_{best,g} + r(U_b - L_b)(z_g - 0.5)$$

where $r \in (0, 1)$ is the chaotic search radius. U_b and L_b are the upper and lower bound vectors of $\bar{z}_{i,g}$ ($\forall i \in 1, 2, \dots, N$), respectively. z_g is a chaotic variable generated at iteration g . If the value of $\bar{z}_{best',g}$ locates out of the search space, it is reset to the closest boundary value. After each execution of CLS, if the fitness of new individual $\bar{z}_{best',g}$ is better than the current global best fitness value $\bar{z}_{best,g}$, $\bar{z}_{best',g}$ replaces it for the next iteration (it is called a successful CLS). The search radius is narrowed at each generation as follows: $r = \rho r$, where $\rho = 0.988$ is a shrinking parameter.

In memory-selectively CJADE-M, one chaotic map is selected from J chaotic maps according to the learned probability from a success rate and a failure rate which are calculated from success and failure memories, respectively. A learning iteration count L is preset. In these learning iterations, chaotic maps are randomly chosen with an equal probability $1/J$.

Table 2

Details of data sets. Abrev: Abbreviation. Class: Number of classes. Att: Number of attributes. Case: Number of cases. θ s: Number of parameters to be optimized.

Data	Abrev	Class	Att	Case	θ s	Data	Abrev	Class	Att	Case	θ s
australian	aust	2	15	690	130	hepatitis	hepa	2	20	80	162
chess	ches	2	37	3296	290	lymphography	lymp	4	19	148	1220
cleveland	clev	2	12	296	1005	Mofn-3-10	mofn	2	11	1324	78
corral	corr	2	7	128	46	pima	pima	2	9	768	102
crx	crxD	2	16	653	848	segment	segm	7	20	2310	2548
diabetes	diab	2	9	768	102	Soybean-large	soyb	19	36	316	5265
flare	flar	8	11	1389	276	Tic-tac-toe	tic-	2	10	958	152
german	germ	2	21	1000	866	vehicle	vehi	4	19	958	1152
glass2	glas	2	10	163	1038	vote	vote	2	18	436	278
heart	hear	2	14	270	118	Waveform-21	wave	3	22	301	3186

while the results of CLS are stored in the success and failure memories. After L learning iterations, just the last L number of successes or failures are recorded. After those learning iterations, each chaotic map has its new chosen probability. The chaotic map with a higher success rate has a higher chance to generate new individuals via a roulette wheel selection method.

3.2. Learning the BN parameters by DE

Being $\mathbf{X}$, a set of discrete random variables of dimension D , each random variable $X_i \in \mathbf{X}$ is associated with a Conditional Probability Table (CPT). An individual $\tilde{z}_{i,g}$, as defined in Section 3 consists of a random variables vector of CPTs in a BN. That is, the estimated CPTs of an individual i at generation g is denoted by $\Theta_{i,g}$. So, an individual is defined as a vector consisting of D CPTs:

$$\tilde{z}_{i,g} = (\Theta_{i,1,g}, \Theta_{i,2,g}, \dots, \Theta_{i,D,g})$$

A particular CPT, e.g., attribute X_a given its parent X_b , corresponds to the second CPT ($D = 2$) of first individual ($i = 1$) at fifth generation ($g = 5$), that CPT is a two entry table given by $\Theta_{1,2,5} = (\theta_{1,1}, \dots, \theta_{1,b}, \dots, \theta_{a,1}, \dots, \theta_{a,b})$, where $\theta_{a,b} \in [0, 1]$ denotes a probability value for a particular state given a parent instantiation. The size of CPT depends on the number of states of the attribute, $s_a = |X_a|$, the number of its parents, $p_a = |\Pi_{X_a}|$, and the number of its parent states $s_{ab} = \sum_{b \in \Pi_{X_a}} |X_b|$, in the following way: $|CPT| = s_a \cdot (s_{ab})^{p_a}$.

A CPT is generated based on the constraint that the sum of probabilities for different states of the random variable should be equal to 1 for a parent instantiation.

Two repairs were applied to satisfy the constraints for $\theta_{a,b} \in [0, 1]$:

$$\theta'_{a,b} = \begin{cases} |\theta_{a,b}| \bmod 1 & \text{if } \theta_{a,b} < 0 \\ 1 - (\theta_{a,b} \bmod 1) & \text{if } \theta_{a,b} > 1 \end{cases}$$

and to keep the sum of row vectors equal to 1:

$$\theta'_{a,b} = \theta_{a,b} / \sum_{b=1}^n \theta_{a,b}.$$

4. Experimental setup

In this section we describe the experiments conducted to compare the performance of our proposed method based on DE with state-of-art discriminative parameter learning called WANBIA in terms of CLL, as well as a comparison in terms of accuracy between generative parameter learning algorithms and other common classifiers. The performance is analyzed on 20 data sets described in Table 2. Records with missing values were deleted and all the data sets were discretized using CAIM [23]. The resulting data sets are in the public domain¹.

In order to determine which DE variants provide a better performance, seven commonly used mutation strategies (rand/1, rand/2, best/1, best2, current-to-rand/1, current-to-best/1 and rand-to-best/1) and two crossover schemes (binomial and exponential) were adopted. JADE without archive, JADE with archive, L-SHADE, and CJADE-M were also included to consider recent variants that have shown a competitive performance.

The performance obtained by the 18 DE variants was compared based on optimized CLL values. To identify the conditions under which different parameter learning algorithms derive better benefits from the discriminative approach, each parameter learning algorithm was applied to the structures learned with the algorithms described in Section 2: Naive Bayes, TAN, and HCS.

¹ <https://sites.google.com/view/vaguilarar/thesis>.

Table 3

Parameter values of DE variants. **F**: Scale factor. **CR**: Crossover rate. **K**: Scale factor. **c**: Rate of parameter adaptation. **p**: Greediness of the mutation strategy. **|A|**: Size of external Archive. **r**: Search radius. **L**: Learning steps. **ρ** : shrinking parameter of search radius.

DE algorithm	F	CR	K	c	p	A	r	L	ρ
rand/1/bin	0.5	0.7							
rand/1/exp	0.5	0.7							
best/1/bin	0.6	0.3							
best/1/exp	0.5	0.8							
best/2/bin	0.5	0.4							
best/2/exp	0.2	0.9							
current-to-best/1/bin	0.7	0.3	F						
current-to-best/1/exp	0.5	0.9	F						
current-to-rand/1/bin	0.6	0.4	F						
current-to-rand/1/exp	0.7	0.9	F						
rand/2/bin	0.3	0.4							
rand/2/exp	0.5	0.8							
rand-to-best/1/bin	0.5	0.4	F						
rand-to-best/1/exp	0.7	0.9	F						
JADE with Archive				0.05	0.05	NP			
JADE without Archive				0.05	0.05	$\emptyset$			
LSHADE				0.05	0.05	NP_g			
CJADE-M				0.05	0.05	NP	0.05	50	0.988

Table 4

Average Ranks for CLL optimization of Differential Evolution variants and WANBIA algorithm in Friedman Test. Values in boldface indicate the best ranked algorithm on each structure.

Algorithm	naïve Bayes		TAN		HCS	
	Average Rank	Final Rank	Average Rank	Final Rank	Average Rank	Final Rank
best/1/bin	4.75	4	5.60	5	6.30	6
best/1/exp	12.75	13	15.05	15	13.35	14
best/2/bin	8.65	9	8.40	9.5	8.30	8
best/2/exp	7.55	7	11.40	12	9.95	11
CJADE-M	3.55	2	3.45	2	4.85	3
current-to-best/1/bin	5.75	6	5.50	4	4.90	4
current-to-best/1/exp	10.50	12	12.95	13	11.15	13
current-to-rand/1/bin	8.25	8	6.80	7	7.20	7
current-to-rand/1/exp	15.60	16	14.70	14	14.80	15
JADE with Archive	3.10	1	2.65	1	3.90	1
JADE without Archive	5.15	5	5.95	6	5.75	5
L-SHADE	4.50	3	3.70	3	4.00	2
rand/1/bin	9.25	10	8.40	9.5	8.45	9
rand/1/exp	15.25	15	16.80	18	15.85	17
rand/2/bin	10.05	11	8.15	8	9.20	10
rand/2/exp	15.70	17	16.05	17	15.50	16
rand-to-best/1/bin	17.60	18	15.85	16	18.75	19
rand-to-best/1/exp	18.20	19	17.90	19	17.35	18
WANBIA	13.85	14	10.70	11	10.45	12

For each DE variant, 31 independent runs were executed on each data set. The DE parameter values used in each variant are detailed in Table 3. Such values were adopted from the specialized literature [19,21,22] and by further experimentation. Regardless of the characteristics of the data set, the initial population NP was established at 100, however, to take into account the different difficulty of the data sets, the number of generations considers the number of attributes multiplied by 15. For CJADE-M, the number of learning iterations L is equal to 50 generations and chaotic search radius r is set to 0.05.

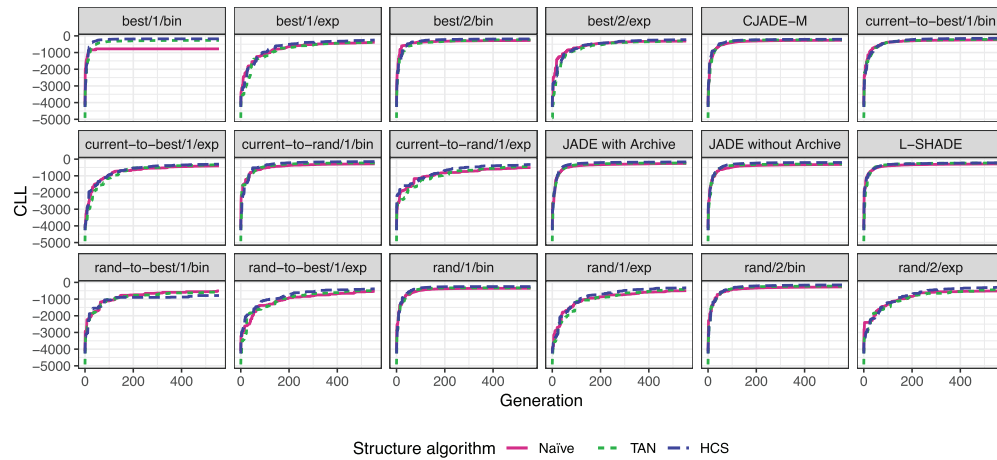
To assess the performance of each algorithm, the Friedman test is adopted as a statistical analysis method to rank each one according to their CLL and accuracy results. It is used to test the performance among different groups. The group with better performance gets a lower rank in each data set. The worst result found in the 31 independent runs is used, this to demonstrate the robustness of the proposed metaheuristic solutions.

In terms of CLL, DE variants and WANBIA are compared. For accuracy, also generative parameter learning algorithms and other standard learning algorithms as Random Forest [24], Artificial Neural Networks [25] or C5.0 [26] are considered. For BN classifiers, comparisons are made over three different structures previously described: naïve Bayes, TAN and HCS.

Table 5

p -values for CLL optimization of Differential Evolution variants and WANBIA algorithm where JADE with Archive is the control algorithm. Values in boldface indicate non-significant differences ($p > .05$).

Algorithm	naïve Bayes				TAN				HCS			
	p -value	p_{Bonf}	p_{Holm}	$p_{Hochberg}$	p -value	p_{Bonf}	p_{Holm}	$p_{Hochberg}$	p -value	p_{Bonf}	p_{Holm}	$p_{Hochberg}$
best/1/bin	0.354	1.000	1.000	0.800	0.097	1.000	0.389	0.328	0.177	1.000	0.887	0.887
best/1/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
best/2/bin	0.002	0.033	0.015	0.015	0.001	0.022	0.011	0.010	0.013	0.241	0.094	0.094
best/2/exp	0.012	0.223	0.074	0.074	0.000	0.000	0.000	0.000	0.001	0.012	0.007	0.007
CJADE-M	0.800	1.000	1.000	0.800	0.653	1.000	1.000	0.653	0.593	1.000	1.000	0.955
current-to-best/1/bin	0.136	1.000	0.682	0.682	0.109	1.000	0.389	0.328	0.574	1.000	1.000	0.955
current-to-best/1/exp	0.000	0.001	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.001	0.001	0.001
current-to-rand/1/bin	0.004	0.068	0.027	0.027	0.020	0.355	0.118	0.118	0.064	1.000	0.382	0.382
current-to-rand/1/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
JADE without Archive	0.249	1.000	0.997	0.800	0.064	1.000	0.318	0.318	0.299	1.000	1.000	0.955
L-SHADE	0.431	1.000	1.000	0.800	0.555	1.000	1.000	0.653	0.955	1.000	1.000	0.955
rand/1/bin	0.001	0.010	0.005	0.005	0.001	0.022	0.011	0.010	0.011	0.190	0.084	0.084
rand/1/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
rand/2/bin	0.000	0.002	0.001	0.001	0.002	0.036	0.014	0.014	0.003	0.052	0.026	0.026
rand/2/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
rand-to-best/1/bin	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
rand-to-best/1/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
WANBIA	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.004	0.003	0.003


Fig. 2. Convergence graph of DE variants on chess data set over three different structures (naïve Bayes, TAN and HCS).

5. Results

This section presents the performance assessment among the 18 DE variants and the WANBIA algorithm, which is the state-of-the-art heuristic approach on Discriminative Learning of parameters. This comparison is in terms of solution quality.

Table 4 summarizes the average rank that each algorithm obtains in CLL optimization over the 20 selected datasets, listing independently the results of the used structure learning algorithms. Based on this table, JADE with Archive gets the top rank on all datasets and structures, obtaining the best performance. L-SHADE and CJADE-M, which are based on JADE with adaptation parameter, also obtain competitive results. The local search of CJADE-M does not seem to have any effects in performance. WANBIA, the state-of-the-art heuristic approach on Discriminative Learning of parameters, was not the most competitive among the compared approaches.

Table 5 shows unadjusted p -values and adjusted p -values of the DE variants and the WANBIA algorithm, where JADE with Archive is the control algorithm on naïve Bayes, TAN and HCS structures. The adjusted p -value is a relatively conservative value used to eliminate the family wise error rate. These p -values smaller than 0.05 indicate a significant difference between JADE with Archive and the other compared algorithms. According to Table 5, it can be concluded that, on the three structures, JADE with Archive significantly outperforms WANBIA and most of the “rand”-based mutation strategies. On the

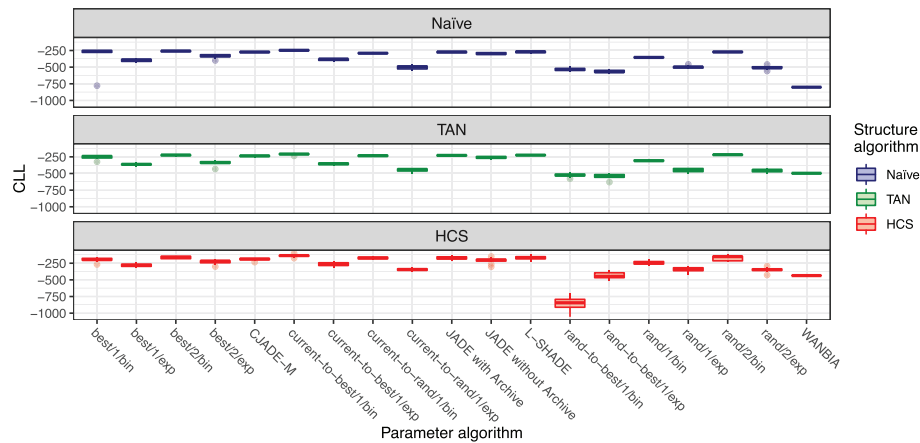


Fig. 3. Box-and-whisker diagrams for solution distribution of DE variants and WANBIA on Chess data set over three different structures (naïve Bayes, TAN and HCS).

Table 6

Average Ranks for Cross Validation accuracy of Differential Evolution variants, WANBIA and generative learning algorithms in Friedman Test. Values in boldface indicate the best ranked algorithm on each structure.

Algorithm	naïve Bayes		TAN		HCS	
	Average Rank	Final Rank	Average Rank	Final Rank	Average Rank	Final Rank
best/1/bin	7.75	6	8.35	7	8.13	7
best/1/exp	13.13	14	16.08	17	15.03	17
best/2/bin	10.40	11	9.78	9	9.23	8
best/2/exp	9.30	8	13.48	15	12.68	15.5
CJADE-M	6.10	1.5	5.40	2	6.13	2
current-to-best/1/bin	7.28	5	7.45	6	6.28	3
current-to-best/1/exp	12.33	13	15.00	16	12.68	15.5
current-to-rand/1/bin	7.20	4	7.03	5	7.95	6
current-to-rand/1/exp	13.65	16	16.45	19	15.08	18
JADE with Archive	6.35	3	5.50	3	6.63	4
JADE without Archive	7.88	7	6.45	4	7.53	5
L-SHADE	6.10	1.5	4.78	1	5.65	1
rand/1/bin	10.38	10	10.15	11	10.05	9
rand/1/exp	15.20	17	18.95	22	16.88	19
rand/2/bin	9.60	9	9.58	8	11.08	11
rand/2/exp	15.83	18	18.13	20	17.15	20
rand-to-best/1/bin	16.73	19	16.28	18	20.78	22
rand-to-best/1/exp	17.70	21	18.43	21	18.40	21
AWNB	13.35	15	12.73	14	11.53	13
Bayesian	18.03	22	10.73	12	12.03	14
ML	16.90	20	9.98	10	11.45	12
WANBIA	11.85	12	12.35	13	10.73	10

other hand, JADE with Archive behaves similar that CJADE-M, JADE without Archive, L-SHADE, best/1/bin and, current-to-best/1/bin. The overall results then suggest that: (1) the adaptation parameter strategy seems to be an effective way to enhance the performance, and (2) including the solutions in the mutation variant improves the observed performance.

To intuitively show the convergence of each DE variant, Fig. 2 includes the plots of the average median CLL solution on the chess data set. The vertical axis is the average median CLL solution and the horizontal axis is the number of generation. As suggested before, Fig. 2 confirms that those mutation strategies where the best solutions in the current population are considered, have faster convergence than other variants and provide the lowest average median CLL solution. Their corresponding box-and-whisker diagrams are plotted in Fig. 3. From it, most DE variants show both, a good and lower solution distribution, i.e., a better search ability than that provided by WANBIA.

The statistical results above mentioned show that the DE approach is a promising and competitive method for Discriminative Learning of BN parameters and significantly outperforms the conventional WANBIA algorithm. From the DE variants compared, those with adaptation parameter schemes and mutation strategies where the best solutions are considered emerge as best options.

Table 7

p -values for Cross Validation accuracy of Differential Evolution variants, WANBIA and generative learning algorithms where L-SHADE is the control algorithm. Values in boldface indicate non-significant differences ($p > .05$).

Algorithm	naïve Bayes				TAN				HCS			
	p -value	p_{Bonf}	p_{Holm}	$p_{Hochberg}$	p -value	p_{Bonf}	p_{Holm}	$p_{Hochberg}$	p -value	p_{Bonf}	p_{Holm}	$p_{Hochberg}$
best/1/bin	0.422	1.000	1.000	1.000	0.082	1.000	0.490	0.490	0.228	1.000	1.000	0.817
best/1/exp	0.001	0.013	0.008	0.008	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
best/2/bin	0.036	0.761	0.363	0.336	0.015	0.313	0.119	0.119	0.082	1.000	0.572	0.572
best/2/exp	0.119	1.000	0.834	0.834	0.000	0.000	0.000	0.000	0.001	0.013	0.009	0.009
CJADE-M	1.000	1.000	1.000	1.000	0.761	1.000	1.000	0.761	0.817	1.000	1.000	0.817
current-to-best/1/bin	0.567	1.000	1.000	1.000	0.193	1.000	0.963	0.761	0.761	1.000	1.000	0.817
current-to-best/1/exp	0.002	0.051	0.029	0.029	0.000	0.000	0.000	0.000	0.001	0.013	0.009	0.009
current-to-rand/1/bin	0.592	1.000	1.000	1.000	0.273	1.000	1.000	0.761	0.263	1.000	1.000	0.817
current-to-rand/1/exp	0.000	0.005	0.004	0.004	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
JADE with Archive	0.903	1.000	1.000	1.000	0.724	1.000	1.000	0.761	0.635	1.000	1.000	0.817
JADE without Archive	0.387	1.000	1.000	1.000	0.415	1.000	1.000	0.761	0.361	1.000	1.000	0.817
rand/1/bin	0.037	0.784	0.363	0.336	0.009	0.186	0.089	0.089	0.032	0.675	0.257	0.257
rand/1/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
rand/2/bin	0.088	1.000	0.706	0.706	0.019	0.408	0.136	0.136	0.008	0.173	0.082	0.082
rand/2/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
rand-to-best/1/bin	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
rand-to-best/1/exp	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
AWNB	0.000	0.009	0.006	0.006	0.000	0.002	0.001	0.001	0.004	0.089	0.051	0.051
Bayesian	0.000	0.000	0.000	0.000	0.004	0.079	0.041	0.041	0.002	0.040	0.025	0.025
ML	0.000	0.000	0.000	0.000	0.011	0.238	0.102	0.102	0.005	0.099	0.052	0.052
WANBIA	0.005	0.107	0.056	0.056	0.000	0.005	0.003	0.003	0.013	0.283	0.121	0.121

Table 8

Average Accuracy, Ranks and p -values in Friedman Test for accuracy of top selected DE variants, WANBIA, generative learning algorithms and other commonly used classifiers. Random Forest is the control algorithm. Values in boldface indicate non-significant differences ($p > .05$).

Algorithm	Avg. Accuracy	Avg. Rank	Final Rank	p -value	p_{Bonf}	p_{Holm}	$p_{Hochberg}$
Random Forest	0.9546	2.275	1				
CJADE-M _{TAN}	0.9042	6.125	2	0.0851	1.0000	0.1546	0.0851
L-SHADE _{TAN}	0.9076	6.225	3	0.0773	1.0000	0.1546	0.0851
JADE with A _{TAN}	0.9049	6.900	4	0.0386	0.8879	0.1158	0.0851
ANN	0.8710	9.125	5	0.0022	0.0503	0.0088	0.0088
ML _{TAN}	0.8937	9.800	6	0.0008	0.0176	0.0038	0.0038
L-SHADE _{HCS}	0.8984	10.200	7	0.0004	0.0091	0.0024	0.0024
C5.0	0.8985	10.400	8	0.0003	0.0064	0.0020	0.0020
CJADE-M _{HCS}	0.8937	10.475	9	0.0002	0.0056	0.0020	0.0020
Bayesian _{TAN}	0.8903	10.850	10	0.0001	0.0029	0.0011	0.0011
JADE with A _{HCS}	0.8949	11.525	11	0.0000	0.0008	0.0004	0.0004
CJADE-M _{NB}	0.8929	12.500	12	0.0000	0.0001	0.0001	0.0001
JADE with A _{NB}	0.8920	12.675	13	0.0000	0.0001	0.0000	0.0000
L-SHADE _{NB}	0.8917	12.725	14	0.0000	0.0001	0.0000	0.0000
WANBIA _{TAN}	0.8852	12.925	15	0.0000	0.0000	0.0000	0.0000
AWNB _{TAN}	0.8770	13.250	16	0.0000	0.0000	0.0000	0.0000
ML _{HCS}	0.8775	14.325	17	0.0000	0.0000	0.0000	0.0000
WANBIA _{HCS}	0.8718	14.875	18	0.0000	0.0000	0.0000	0.0000
AWNB _{HCS}	0.8621	15.775	19	0.0000	0.0000	0.0000	0.0000
Bayesian _{HCS}	0.8620	16.000	20	0.0000	0.0000	0.0000	0.0000
WANBIA _{NB}	0.8544	17.950	21	0.0000	0.0000	0.0000	0.0000
AWNB _{NB}	0.8333	19.225	22	0.0000	0.0000	0.0000	0.0000
ML _{NB}	0.8407	21.400	23	0.0000	0.0000	0.0000	0.0000
Bayesian _{NB}	0.8339	22.475	24	0.0000	0.0000	0.0000	0.0000

Table 6 summarizes the average rank that each algorithm obtains in accuracy over the 20 selected datasets, listing the results independently by the adopted structure learning algorithm. L-SHADE gets the top rank on all datasets and structures, indicating that it has the best performance. CJADE-M and JADE with Archive also obtain some competitive results, but below CJADE-M. The DE variants with random solutions as base vector provided poor results. The performance of the Generative Learning algorithms is not very competitive in general, but it is worse when the assumptions of independence between the parameters is strong, as in the naïve Bayes structure. The performance provided by WANBIA, although it includes discriminative learning, is poor.

Table 7 shows unadjusted p -values and adjusted p -values for accuracy of DE variants, WANBIA algorithm and three generative algorithms, where L-SHADE is the control algorithm on naïve Bayes, TAN and HCS structures. These p -values

smaller than 0.05 indicate the significant difference between L-SHADE and other compared algorithms. According to Table 7, it can be concluded that on the three structures, L-SHADE performs similarly with respect to C-JADE-M and JADE.

As expected, L-SHADE outperforms Generative Learning of parameters algorithms when structure encodes strong independence assumptions among attributes. WANBIA has competitive results in naïve Bayes and HCS structures against L-SHADE but it is not the case in TAN structure. Those results suggest that: (1) as in previous experiments, the bias provided by those best solutions in the mutation operator coupled with adaptation parameter schemes provide very competitive results for the classification task even when strong independence assumptions are made, (2) L-SHADE, C-JADE-M and JADE perform well even when optimizing highly connected networks and the search space becomes complex as in the TAN structure.

Table 8 summarizes the average accuracy, average rank, and p -values that the compared algorithms obtained in accuracy over 20 datasets. Based on this table, it can be seen that the proposed methods obtain reasonable good results compared with other standard learning algorithms as Random Forest [24], Artificial Neural Networks [25] or C5.0 [26], especially when the network structure is good. We must remember that the worst case found by the DE variants is the compared one. We must also emphasize that, even if the network is not good, the proposed approach guarantees the best result compared to other discriminative or generative parameter learning algorithms in BN.

6. Conclusions

Discriminative Learning the parameters of a BN, an open NP-hard problem, has been tackled using different representative DE variants. The paper focused on BN because it is typically easier to acquire meaningful models with them, given their capacities to express prior knowledge and also the existing algorithms for learning their structure from data. Therefore, BN might be preferred to explicitly represent and exploit relations among attributes.

The meta-heuristic approach in Discriminative Learning, especially in the context of BN, is still in its infancy stage. This paper might help to introduce these ideas in the Machine Learning community, showing that the approach can work effectively. This would lead to the generation of classifiers with low bias that minimize the classification error.

Based on the results, it was observed that the mutation strategies with randomly chosen base vectors present a slower convergence than those where the best solutions are considered instead. In this regard, although DE variants like best/1/bin or current-to-best/1/bin have a fast convergence and are competitive in terms of CLL, the calibration of the F and CR parameters is not trivial and requires prior experimentation. For this reason, the DE variants based on JADE, such as L-SHADE and C-JADE, are especially recommended because of their suitable parameter adaptation mechanisms, which also keeps the user from a careful calibration effort. It should be noted that, although the differences between JADE with Archive and without Archive are not statistically significant, the inclusion of the external JADE file has a positive effect on search by maintaining diversity, so including such file is recommended.

The chaotic local search of C-JADE did not have positive effects on the optimization process. However, there is the possibility that by performing a different calibration from that proposed by its authors [22], better results may be obtained. Compared with JADE with Archive, the linear population reduction mechanism of L-SHADE does not significantly help to provide better results in terms of CLL. However, L-SHADE works with smaller populations mainly later in the search. Therefore, it is recommended over any other DE variant if the cost is an issue.

The variability in the distribution of the solutions is low for most of DE variants, which indicates that the proposed approach is robust and even in the worst case there is a high quality solution when compared against WANBIA. Regarding the CLL optimization, those JADE variants with parameter adaptation mechanisms were better in the three explored structures. In terms of classification, L-SHADE was also superior, mainly in the TAN structure. This might suggest that the DE variants, specifically those based on JADE, have a higher search performance than WANBIA, regardless of the complexity of the search space.

Finally, our results are consistent with the analysis of previous approaches for Discriminative Learning the parameters of a BN. In particular, the benefit of the approach is considerable for simple network structures, which are not optimized for classification, computing in this case more accurate classifiers than Generative Learning.

We show that the DE approach works effectively over a variety of situations: when dealing with structures that range from trivial (naïve Bayes), through less-trivial (HCS), to complex ones (TAN). We also show that the DE approach works well in essentially every standard situation, we see then that the DE approach is at least as good, and often better, than the state-of-art method for Discriminative Learning called WANBIA.

Future work includes the application of meta-heuristic approaches for Discriminative Learning now for both structure and parameters. A detailed and more thorough experimentation on C-JADE implementation is required to include the benefits from the local search. Conducting an extended analysis with other data sets, structures and DE variants, is considered with the aim of identifying the conditions under which different strategies achieve better performance.

References

- [1] N. Friedman, D. Geiger, M. Goldszmidt, Bayesian network classifiers, *Mach. Learn.* 29 (1997) 131–163.
- [2] N. Zaidi, G. Webb, M. Carman, F. Petitjean, W. Buntine, M. Hynes, H. Sterck, Efficient parameter learning of bayesian network classifiers, *Mach. Learn.* 106 (2017) 1289–1329.



- [3] J. Su, H. Zhang, C. Ling, S. Matwin, Discriminative parameter learning for bayesian networks, in: In Proc. of the 25th international conference on Machine learning, ACM, 2008, pp. 1016–1023.
- [4] K. Price, R. Storn, Minimizing the real functions of the iccc'96 contest by differential evolution, in: in Proc. of IEEE C. Evol. Computat., 1996, pp. 842–844.
- [5] O. Tarkhaneh, H. Shen, An adaptive differential evolution algorithm to optimal multi-level thresholding for MRI brain image segmentation, Expert Syst. Appl. 138 (2019) 112820, doi:10.1016/j.eswa.2019.07.037.
- [6] Bilal, M. Pant, H. Zaheer, L. García-Hernandez, A. Abraham, Differential evolution: a review of more than two decades of research, Eng. Appl. Artif. Intell. 90 (2020) 103479, doi:10.1016/j.engappai.2020.103479.
- [7] M. Minsky, Steps toward artificial intelligence, Proc. IRE 49 (1) (1961) 8–30, doi:10.1109/jrproc.1961.287775.
- [8] E.J. Keogh, M.J. Pazzani, Learning the structure of augmented bayesian classifiers, Int. J. Artif. Intell. Tools 11 (04) (2002) 587–601, doi:10.1142/s0218213002001052.
- [9] S. Russell, P. Norvig, Artificial intelligence: A Modern approach, 2nd edition, Prentice-Hall, Englewood Cliffs, NJ, 2003.
- [10] M. Scanagatta, A. Salmerón, F. Stella, A survey on bayesian network structure learning from data, Prog. Artif. Intell. 8 (2019) 425–439.
- [11] C. Chow, C. Liu, Approximating discrete probability distributions with dependence trees, IEEE Trans. Inf. Theory 14 (1968) 462–467.
- [12] P. Sundararajan, O. Mengshoel, A genetic algorithm for learning parameters in bayesian networks using expectation maximization, in: in Proc. of the 8th International Conference on Probabilistic Graphical Models, 2016, pp. 511–522.
- [13] R. Greiner, W. Zhou, Structural extension to logistic regression: Discriminative parameter learning of belief net classifiers, in: AAAI/IAAI, 2002, pp. 167–173.
- [14] B. Shen, X. Su, R. Greiner, P. Musilek, C. Cheng, Discriminative parameter learning of general bayesian network classifiers, in: in Proc. Int. C. Tools Art., 2003, pp. 296–305.
- [15] R. Greiner, X. Su, B. Shen, W. Zhou, Structural extension to logistic regression: discriminative parameter learning of belief net classifiers, Mach. Learn. 59 (2005) 297–322.
- [16] R. Greiner, A.J. Grove, D. Schuurmans, Learning bayesian nets that perform well, in: Proceedings of the Thirteenth Conference on Uncertainty in Artificial Intelligence, in: UAI'97, Morgan Kaufmann Publishers Inc., San Francisco, CA, USA, 1997, pp. 198–207.
- [17] F. Pernkopf, M. Wöhlmayr, On discriminative parameter learning of bayesian network classifiers, in: in Proc. ECML PKDD, 2009, pp. 221–237.
- [18] N.A. Zaidi, J. Cerquides, M.J. Carman, G.I. Webb, Alleviating naive bayes attribute independence assumption by attribute weighting, J. Mach. Learn. Res. 14 (2013) 1947–1988.
- [19] J. Zhang, A. Sanderson, Jade: adaptive differential evolution with optional external archive, IEEE Trans. Evol. Comput. 13 (2009) 945–958.
- [20] E. Mezura-Montes, J. Velázquez-Reyes, C.A.C. Coello, A comparative study of differential evolution variants for global optimization, in: Proceedings of the 8th annual conference on Genetic and evolutionary computation - GECCO '06, ACM Press, 2006, pp. 485–492, doi:10.1145/1143997.1144086.
- [21] R. Tanabe, A. Fukunaga, Improving the search performance of shade using linear population size reduction, in: in Proc. of IEEE C. Evol. Computat., 2014, pp. 1658–1665.
- [22] S. Gao, Y. Yu, Y. Wang, J. Wang, J. Cheng, M. Zhou, Chaotic local search-based differential evolution algorithms for optimization, IEEE Transactions on Systems, Man, and Cybernetics: Systems (2020) 1–14, doi:10.1109/tsmc.2019.2956121.
- [23] L. Kurgan, K. Cios, CAIM Discretization algorithm, IEEE Trans. Knowl. Data Eng. 16 (2) (2004) 145–153, doi:10.1109/tkde.2004.1269594.
- [24] L. Breiman, Random forests, Mach. Learn. 45 (1) (2001) 5–32, doi:10.1023/a:1010933404324.
- [25] B.D. Ripley, Pattern recognition and neural networks, Cambridge University Press, 1996, doi:10.1017/cbo9780511812651.
- [26] J.R. Quinlan, Improved use of continuous attributes in c4.5, Journal of Artificial Intelligence Research 4 (1996) 77–90, doi:10.1613/jair.279.



Original software publication

dp1bnDE: An R package for discriminative parameter learning of Bayesian Networks by Differential Evolution

Alejandro Platas-López^{a,c,*}, Alejandro Guerra-Hernández^a, Francisco Grimaldo^b, Nicandro Cruz-Ramírez^a, Efrén Mezura-Montes^a, Marcela Quiroz-Castellanos^a^a Instituto de Investigaciones en Inteligencia Artificial, Universidad Veracruzana, Xalapa, Ver., 91097, Mexico^b Departament d'Informàtica, Universitat de València, Burjassot, 46100, Spain^c Extensión Académica, Universidad Anahuac Online, Ciudad de México, 01840, Mexico

ARTICLE INFO

Article history:

Received 7 May 2023

Received in revised form 6 June 2023

Accepted 12 June 2023

Keywords:

Bayesian Networks
Discriminative Learning
Differential Evolution
Parameter optimization

ABSTRACT

The dp1bnDE R package is a novel tool that implements Differential Evolution strategies for training Bayesian Network parameters using Discriminative Learning. Focusing on optimizing the Conditional Log-Likelihood rather than the log-likelihood, dp1bnDE enhances the performance of Bayesian Networks models in various applications. The package offers four main functions (DErand, DEbest, jade, and lshade) that implement different DE variants, providing users with a versatile and efficient approach to Bayesian Network parameter learning. dp1bnDE has the potential to impact data-driven industries by improving predictive capabilities and decision-making processes in fields such as healthcare, finance, and supply chain management. The package and its code are made freely available. © 2023 The Author(s). Published by Elsevier B.V. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

Code metadata

Current code version	v0.1.1
Permanent link to code/repository used for this code version	https://github.com/ElsevierSoftwareX/SOFTX-D-23-00297
Code Ocean compute capsule	NA
Legal Code License	GPL (≥ 2)
Code versioning system used	git
Software code languages, tools, and services used	R
Compilation requirements, operating environments & dependencies	R (version ≥ 3.2.0), R package: bnclassify (version ≥ 0.4.5)
If available Link to developer documentation/manual	https://github.com/alexplatasl/dp1bnDE
Support email for questions	alejandro.plataslo@anahuac.mx

1. Motivation and significance

Bayesian Networks (BN) have gained significant attention in various scientific domains due to their ability to model complex relationships among variables in a probabilistic manner. Traditionally, BN are trained using Generative Learning approaches that optimize the log-likelihood of the data. However, these approaches may not produce optimal classifiers when the primary goal is to minimize classification error, as they focus on modeling

the joint probability distribution of variables, which might not be directly related to the classification task. Discriminative Learning, which optimizes the Conditional Log-Likelihood, has emerged as a promising alternative that prioritizes classification accuracy by directly modeling the posterior probability of the target class. This work was motivated by the need for an efficient and robust method for experimenting with the Discriminative Learning of parameters in BN using meta-heuristic approaches, such as Differential Evolution (DE). Efficiency refers to finding high-quality solutions with fewer iterations, while robustness indicates consistent performance across various problem instances.

The proposed software package introduces a novel DE-based approach for discriminative parameter learning in BN. It addresses the limitations of existing methods by offering robustness and adaptability in the parameter learning process. The software is important as it provides researchers and practitioners with an

* Corresponding author at: Instituto de Investigaciones en Inteligencia Artificial, Universidad Veracruzana, Xalapa, Ver., 91097, Mexico.

E-mail addresses: alejandro.plataslo@anahuac.mx (Alejandro Platas-López), aguerra@uv.mx (Alejandro Guerra-Hernández), francisco.grimaldo@uv.es (Francisco Grimaldo), ncruz@uv.mx (Nicandro Cruz-Ramírez), emezura@uv.mx (Efrén Mezura-Montes), maquiroz@uv.mx (Marcela Quiroz-Castellanos).

<https://doi.org/10.1016/j.softx.2023.101442>

2352-7110/© 2023 The Author(s). Published by Elsevier B.V. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

easy-to-use tool for BN parameter learning that can produce more accurate classifiers than generative learning, particularly when strong independence assumptions are present among attributes. Adaptability, in this context, refers to the method's capacity to perform well across a diverse set of BN structures and data distributions, adjusting its search strategy accordingly.

The software has the potential to significantly impact various fields where BN are employed, such as bioinformatics [1,2], finance [3,4], and environmental sciences [5,6], by improving the accuracy of classification tasks. Among the range of available discriminative learning tools, the DE-based approach utilized by *dp1bnDE* has been proven to be notably effective. Its performance has been comprehensively evaluated and favorably compared not only to the state-of-the-art method for discriminative learning, known as Weighting Attributes to Alleviate Naïve Bayes' Independence Assumption (WANBIA) [7], but also to other prevalent algorithms used in Bayesian network classifiers. *dp1bnDE* is more effective than WANBIA in the sense that it offers a versatile and efficient approach to BN parameter learning through the implementation of different DE variants, which allows for better adaptation to diverse problem characteristics and data distributions. We anticipate further contributions to scientific discovery as more researchers adopt and utilize this software for their discriminative learning tasks.

The *dp1bnDE* package is designed to be used in the R statistical programming environment, providing an accessible and user-friendly interface for researchers and practitioners working with BN. Users can employ the package to apply Discriminative Learning of parameters in BN using various DE strategies. To use the package, users need to install it from CRAN or GitHub and load it into their R environment. Then, they can choose a DE variant (*DErand*, *DEbest*, *jade*, or *lshade*) based on their preference or prior knowledge about the problem at hand, and input their dataset, network structure, and other relevant parameters to perform the parameter learning process. The package outputs the learned parameters, which can be further analyzed or used for predictions and other tasks. As more research is conducted, users may gain insights into which DE variant performs better in specific cases.

The *dp1bnDE* package builds upon existing research on BN and DE algorithms. BN have been extensively studied and applied in various domains, including artificial intelligence, machine learning, and statistical modeling [8]. Discriminative Learning of BN parameters, which focuses on minimizing classification errors, has been explored by [9].

In terms of optimization algorithms, DE has been widely used for global optimization in continuous spaces [10]. The *dp1bnDE* package incorporates several DE variants, including *DE/rand*, *DE/best*, Adaptive Differential Evolution with Optional External Archive (JADE) [11], and Success History Based Adaptive Differential Evolution with Linear Population Size Reduction (LSHADE) [12], each of which has been successfully applied to various optimization problems.

Other software packages related to BN and optimization include *bnlearn* [13] and *deal* [14] for BN structure learning, *DEoptim* [15] for general DE optimization in R, and *pcaIsg* [16] which focuses on causal structure learning and estimation of causal effects from observational data. However, the *dp1bnDE* package uniquely combines Discriminative Learning of BN parameters with various DE strategies. This integration aims to provide an alternative and useful approach for researchers and practitioners working with BN, potentially enhancing their analysis and results.

2. Software description

The *dp1bnDE* software package is designed for parameter learning in Bayesian Networks (BNs) using Differential Evolution (DE) optimization techniques. By providing an accessible interface within the R environment, the software streamlines the BN parameter learning process, making it user-friendly for a diverse audience. Key components of the software vocabulary include:

Bayesian Networks: Probabilistic graphical models representing relationships among variables, enabling inference, prediction, and decision-making.

Discriminative Learning: An approach that optimizes the Conditional Log-Likelihood of BN parameters to improve classification accuracy.

Differential Evolution: A meta-heuristic optimization algorithm that iteratively updates BN parameters, enhancing the model's performance on input data.

DE Variants: Four available DE strategies (*DErand*, *DEbest*, *jade*, and *lshade*) for optimizing BN parameters, offering adaptability and robustness in the learning process.

Network Structure: A user-defined or data-derived structure representing the relationships among variables in a BN.

In a typical workflow, users input a dataset containing variables of interest and a network structure. The software then employs one of the four DE variants to optimize the BN parameters' Conditional Log-Likelihood. Upon completion of the optimization process, the package outputs a trained BN model suitable for tasks such as classification, prediction, and decision-making.

2.1. Software architecture

The *dp1bnDE* package is designed with a modular architecture, as illustrated in Fig. 1. It consists of four main functions: *DErand*, *DEbest*, *jade*, and *lshade*. Each function represents a distinct DE variant, offering flexibility in choosing the appropriate optimization method for a given problem. The package is implemented in R and can easily integrate with other R libraries and data formats.

2.2. Software functionalities

The primary functionality of the *dp1bnDE* package is to estimate the parameters of BN given a data set and a network structure. To achieve this goal, the package incorporates different variants of the DE algorithm, which enable users to efficiently learn BN parameters.

The major features of the *dp1bnDE* package are as follows:

1. Estimation of BN parameters: Based on the provided data set and network structure, the package allows users to estimate the parameters of the BN using the *DE/rand*, *DE/best*, *JADE*, and *L-SHADE* variants of DE.
2. Flexible configuration: Users can configure the DE settings by selecting the appropriate mutation strategy, number of vector differences, and recombination type, allowing for optimal tuning of the learning process.
3. Performance evaluation: The package allows for assessing the classification performance of the learned BN models, enabling users to compare the results of different DE variants and settings based on the classification performance metric.
4. Reproducible and extendable code: As the package and its source code are made freely available, users can reproduce the results, modify the code, or extend the functionalities to fit their specific needs.

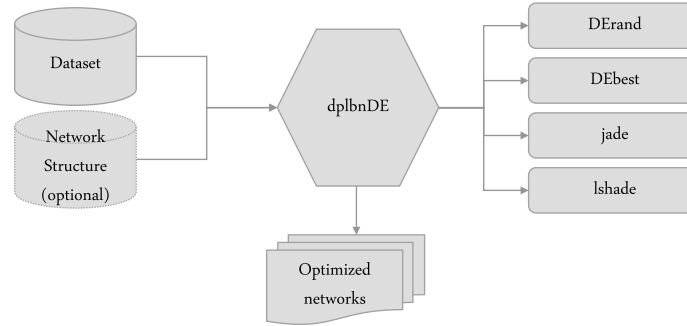


Fig. 1. Modular architecture of the developed dplbnDE package.

Table 1

Summary of common arguments and their descriptions.

Argument	Description
NP	Controls the population size
G	Specifies the number of generations
data	Specifies the data set to be used
class.name	Specifies the name of the class variable in the data set
structure	Specifies the algorithm for learning the network structure
edgelist	Allows for the inclusion of a different structure provided by the user
verbose	Specifies how many generations to print optimization progress on the screen

3. Illustrative examples

This section is divided into two distinct subsections to provide a comprehensive understanding of the capabilities of the dplbnDE package. The first subsection demonstrates the process of classification by learning parameters using different DE variants. This subsection offers a step-by-step illustration of how the package's DE strategies can be employed to improve the performance of BN models. The second subsection presents an example of modifying the structure of the network. This subsection highlights the package's flexibility in accommodating different network configurations and showcases its potential for broader applications. By exploring these two subsections, readers will gain a thorough understanding of the key functionalities and benefits provided by the dplbnDE package.

3.1. Parameter learning

The dplbnDE package, along with its DE variants (DErand, DEbest, jade, and lshade), will be compared in a classification task. The experimental setup consists of the following steps. First, the dplbnDE package is loaded. Next, the car dataset is imported and a seed is set for reproducibility. A train-test split proportion of 0.8 is established in this example, and the number of training examples is calculated. Random indices are generated to create the training and test sets. We would like to note that while this example uses a train-test split (or holdout) sampling strategy, the dplbnDE package is compatible with other strategies, including cross-validation. The choice of sampling strategy can be adjusted according to the needs of the specific analysis. This example, therefore, offers a succinct and comprehensive overview of the data preprocessing steps necessary before moving on to parameter learning and evaluation, adaptable to different sampling strategies.

```

1 # Load packages
2 library(dplbnDE)
3
4 # Load data
5 data(car)
6
7 # Set seed for reproducibility
8 set.seed(1234)
9
10 # Proportion for train-test split
11 proportion <- 0.8
12
13 # Calculate the number of training samples
14 n_train <- round(nrow(car) * proportion)
15
16 # Generate random indices for the training set
17 index_train <- sample(1:nrow(car), n_train)
18
19 # Create the training set using the random indices
20 data_train <- car[index_train, ]
21
22 # Create the test set using the indices not in the
23   training set
24 data_test <- car[-index_train, ]

```

In the dplbnDE package, certain arguments are common to all DE variants implemented. A comprehensive list of these arguments is provided in Table 1.

These arguments provide users with a flexible and customizable approach to DE-based BN parameter learning in dplbnDE. Now, each DE strategy implemented in the package has specific arguments in addition to the common ones. Furthermore, all functions are designed with default parameters that have been carefully chosen based on recommended values from the literature and thorough experimentation on different instances. For DErand and DEbest, these include the crossover rate CR and the scale factor F. The CR argument determines the extent to which the parent influences the generation of offspring, while the F argument scales the impact of the pairs of solutions selected for calculating the mutation value. Additionally, the mutation.pairs argument controls the number of pairs of individuals chosen and can take values of either one or two pairs. Lastly, the crossover argument specifies the type of crossover to be used, with options of binomial and exponential crossover.

In this example, a BN structure is required as input for the application of DE variants. The chosen structure is learned using

the Naïve Bayes algorithm; however, alternative structures can also be constructed using other methods, such as Tree Augmented Network (TAN-CL) or Hill Climbing algorithms. Once the structure is obtained, the DE/rand/1/bin and DE/best/2/exp variants are applied, utilizing 100 generations and a population of 40 individuals, to the BN structure derived from the Naïve Bayes algorithm.

```
1 dpl.rand <- DErand(NP=40, G=100, data = data_train,
  ↪ class.name = "class", F = 0.5, CR = 0.5,
  ↪ mutation.pairs = 1, crossover = "bin", structure =
  ↪ "nb")

1 dpl.best <- DEbest(NP=40, G=100, data = data_train,
  ↪ class.name = "class", F = 0.3, CR = 0.9,
  ↪ mutation.pairs = 2, crossover = "exp", structure =
  ↪ "nb")
```

Subsequently, the jade variant is applied, which has three specific arguments: *c* controls the rate of optimization parameter adaptation, *pB* determines the greediness of the mutation strategy, and *archive* designates when to use an archive of discarded solutions during optimization in order to potentially create the mutation vector, thus fostering diversity. In the provided example, the chosen hyperparameters include a *c* rate of 0.1, *pB* greediness of 0.05, *archive* set to true, and a verbosity level of 10.

```
1 dpl.jade <- jade(NP=40, G=100, data = data_train,
  ↪ class.name = "class", c = 0.1, structure = "nb",
  ↪ pB=0.05, archive = TRUE)
```

Finally, the lshade variant is employed, which is an enhancement of JADE that incorporates a linear reduction of the population size. Consequently, it also includes the *c* and *pB* parameters, and by default, utilizes the archive. In the provided example, an initial population size of 40 is used, which gradually decreases over the course of 100 generations.

```
1 dpl.lshade <- lshade(NP=40, G=100, data = data_train,
  ↪ class.name = "class", c = 0.1, structure = "nb",
  ↪ pB=0.05)
```

For any of the variants, it is possible to inspect a summary of the optimization results from the BN parameter learning process. The summary results can be interpreted as follows:

1. Number of evaluations: The total number of times the objective function (Conditional Log-Likelihood, or CLL) was evaluated during the optimization process.
2. Final population size: The number of candidate solutions in the final population after the optimization process has been completed.
3. Best CLL: The highest (least negative) CLL value achieved by the best candidate solution in the final population, indicating the best fitness.
4. Worst CLL: The lowest (most negative) CLL value among the candidate solutions in the final population, indicating the poorest fitness.

5. Median: The median CLL value of the candidate solutions in the final population, providing a measure of central tendency for the optimization results.
6. Std. Dev.: The standard deviation of the CLL values in the final population, indicating the dispersion or spread of the fitness values.

```
1 print(dpl.lshade)
2 # Number of evaluations:      2660
3 # Final population size:     16
4 #
5 # Summary results of fitness in final population:
6 #
7 # Best CLL:                  -330.3376
8 # Worst CLL:                 -342.5009
9 # Median:                    -338.9794
10 # Std. Dev.:                 3.844851
```

By analyzing these metrics, users can gain a better understanding of the optimization algorithm's performance and the quality of the solutions found during the BN parameter learning process. Similarly, the optimization results of the BN parameters can be visually represented using the *plot* function. The resulting plot is shown in Fig. 2.

```
1 plot(dpl.jade)
```

Next, the classification performance of each method is compared on the test data set. The *predict()* function is used to generate predictions for the test dataset using the best model learned by each DE variant (*dpl.rand\$Best*, *dpl.best\$Best*, *dpl.jade\$Best*, and *dpl.lshade\$Best*). The *accuracy()* function then compares these predictions to the true class labels in the *data_test\$class* column, calculating the classification performance for each DE variant. Finally, the classification performance values for all four DE variants are combined into a single named vector called *accuracy_test*.

```
1 # Calculate the classification performance for each DE
  ↪ variant on the test data
2 # 1. Use the 'predict()' function to make predictions
  ↪ using the best model for each variant
3 # 2. Compare the predictions to the true class labels
  ↪ in the test dataset
4 # 3. Use the 'accuracy()' function to compute the
  ↪ classification performance for each variant

# Classification performance for DE/rand variant
accuracy_rand <- accuracy(predict(dpl.rand$Best,
  ↪ data_test), data_test$class)

# Classification performance for DE/best variant
accuracy_best <- accuracy(predict(dpl.best$Best,
  ↪ data_test), data_test$class)

# Classification performance for JADE variant
accuracy_jade <- accuracy(predict(dpl.jade$Best,
  ↪ data_test), data_test$class)

# Classification performance for L-SHADE variant
accuracy_lshade <- accuracy(predict(dpl.lshade$Best,
  ↪ data_test), data_test$class)
```

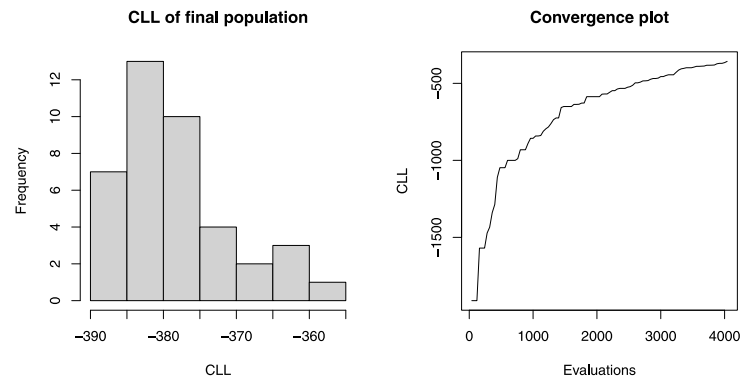


Fig. 2. Histogram of optimized Conditional Log-Likelihood values for the final population and convergence plot of the CLL with jade variant.

```

18 # Combine classification performance values into a
19   ↪ single vector
20 accuracy_test <- c(
21   rand = accuracy_rand,
22   best = accuracy_best,
23   jade = accuracy_jade,
24   lshade = accuracy_lshade
25 )
26 # Print results
27 print(accuracy_test)
28 #      rand      best      jade      lshade
29 #0.8699422 0.8352601 0.8612717 0.8815029

```

```

1 # Create the matrix
2 my_structure <- matrix(c("class", "buying",
3   "class", "maint",
4   "class", "doors",
5   "class", "persons",
6   "class", "lug_boot",
7   "class", "safety",
8   "maint", "buying",
9   "lug_boot", "safety"),
10   nrow = 8, ncol = 2, byrow = TRUE)
11
12 # Set column names
13 colnames(my_structure) <- c("from", "to")
14
15 run.shade = lshade(NP=40, G=100, data = car,
16   ↪ class.name = "class", c = 0.1, pB=0.05, edgelist =
17   ↪ my_structure)
18 plot(run.shade$Best)

```

The \$Best list element in each DE variant object (e.g., dpl.rand\$Best, dpl.best\$Best, dpl.jade\$Best, and dpl.lshade\$Best) holds the optimal parameters found in terms of Conditional Log-Likelihood (CLL). These parameters have been fine-tuned through the DE algorithm to maximize the CLL. By using the \$Best list element, we ensure that the most suitable set of parameters is employed when assessing the classification performance of each DE variant on both the training and test data sets.

3.2. Change structures

The dplbnDE package allows users to provide any BN structure by specifying the relationships between variables through an edge list represented as a two-column matrix. In the provided R code, a custom structure is defined using a matrix, my_structure, which represents the relationships between the nodes in the network. The from and to columns in the matrix denote the directed edges between the nodes. The edge list is then passed as an argument to the lshade function, which is used to learn the parameters of the BN using the custom structure. The edge list is specified using the edgelist parameter. After learning the parameters, the plot function is used to visualize the structure of the BN (see Fig. 3). This feature of dplbnDE enables users to impose their domain knowledge on the network's structure, potentially improving model performance and interpretability.

4. Impact

The dplbnDE package has the potential to substantially impact research and application of BNs, particularly in Discriminative Learning by optimizing the Conditional Log-Likelihood. Offering a convenient and flexible platform for implementing and comparing various DE strategies, it promotes the development of novel optimization techniques tailored to parameter learning. This new research direction not only spurs innovation in the BN field but also contributes to a deeper understanding of the interplay between metaheuristics and parameter learning in other machine learning domains.

The package streamlines the research process by providing an easy-to-use interface for implementing various DE strategies, allowing researchers to focus on core research questions without being burdened by optimization algorithm implementation details. This accelerates research and encourages experimentation. Incorporating multiple DE variants, including DE/rand, DE/best, JADE, and LSHADE, the dplbnDE package enables researchers to swiftly compare and evaluate different optimization strategies for BN parameter learning. This unified platform promotes rigorous evaluation and benchmarking of optimization techniques.

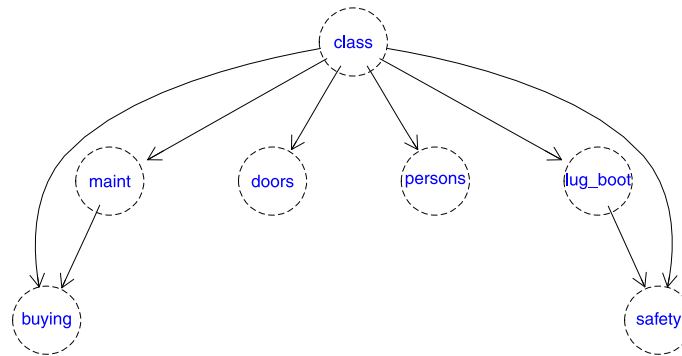


Fig. 3. Custom Bayesian Network: Visualizing the relationships between variables using a user-defined edge list.

Moreover, the `dplbnDE` package offers a high degree of customization, enabling researchers to tailor optimization parameters and BN structures to their specific research questions. This flexibility allows for a more in-depth exploration of the impact of various factors on parameter learning. Utilizing advanced optimization techniques, such as LSHADE's Linear Population Size Reduction (LPSR), the `dplbnDE` package contributes to developing more efficient and accurate parameter learning methods, leading to the discovery of improved BN models and insights into underlying data.

The package simplifies discriminative parameter learning in BN by providing easy-to-use functions and pre-configured optimization strategies, lowering the barrier for researchers and practitioners who may lack extensive experience with optimization techniques. Offering a user-friendly interface and a comprehensive set of DE variants, the package empowers users to explore and apply methods with ease. It also streamlines the process of experimenting with different DE strategies, enabling users to quickly compare and evaluate their performance on a given problem. This efficiency in experimentation fosters a deeper understanding of the strengths and weaknesses of various optimization techniques, ultimately leading to more informed decisions in model selection and parameter tuning.

Furthermore, the customizable nature of the `dplbnDE` package allows users to tailor their analyses to specific research questions or datasets, encouraging more targeted and thorough investigation of discriminative parameter learning. This adaptability facilitates the exploration of new research questions and promotes a more nuanced understanding of the factors that influence BN performance.

While the `dplbnDE` package is relatively new, it offers features that could be of interest to researchers and practitioners working on Bayesian Networks and optimization techniques. Its user-friendly interface and comprehensive set of DE variants make it an attractive tool for those looking to explore and apply discriminative parameter learning in BNs. Additionally, the package's optimization techniques and BN modeling capabilities may attract users from related fields such as machine learning, data mining, and artificial intelligence.

The package's applicability extends to various domains and industries where BNs are employed for decision-making and predictive modeling, such as healthcare [17,18], finance [3,4,19], marketing [20], and environmental sciences [5,6]. Professionals in these industries can benefit from `dplbnDE`'s efficient and accurate discriminative parameter learning methods to improve their BN models and gain deeper insights into their data.

We posit that the `dplbnDE` package could have commercial relevance given its focus on discriminative parameter learning in BN, an area utilized across a wide array of industries for decision-making and predictive modeling. For instance, in the healthcare sector, providers can utilize BNs to model complex relationships between medical conditions, symptoms, and treatments [17]. Potential applications of `dplbnDE` could extend to healthcare, where the optimized discriminative learning of these networks may aid professionals in enhancing the performance of diagnostic and treatment recommendations. This suggests a potential for improved patient outcomes and more efficient resource allocation, subject to further research and validation.

Similarly, in the finance industry, BNs can be used to model relationships between economic indicators, market trends, and investment opportunities [21–23]. The optimization techniques provided by `dplbnDE` can enhance the predictive capabilities of these models, enabling financial institutions to make more informed investment decisions and better manage risk.

Moreover, the optimization techniques offered by `dplbnDE` can potentially be applied to various problems beyond BN, such as supply chain optimization [24], portfolio management [25], and cloud resource allocation [26]. This versatility may attract companies specializing in data-driven decision-making, analytics, or artificial intelligence, leading to further adoption in commercial settings.

While it is too early to determine whether `dplbnDE` will directly lead to the creation of spin-off companies, its functionality and potential applications across different industries suggest that it holds promise for making a meaningful impact in commercial settings. As the package gains recognition and adoption, it may contribute to driving innovation and improving decision-making processes in data-driven industries.

5. Conclusions

In conclusion, the `dplbnDE` package is a significant contribution to the field of BN parameter learning, as it focuses on optimizing the Conditional Log-Likelihood (Discriminative Learning) instead of the log-likelihood (Generative Learning). The choice to optimize the Conditional Log-Likelihood can result in more accurate classifiers, especially when minimizing classification errors is the primary goal. By implementing various DE strategies, `dplbnDE` provides researchers and practitioners with an efficient and versatile approach to train the parameters of BNs, enhancing the performance and the classification error of these models in real-world applications.

The package offers four main functions (DErand, DEbest, jade, and lshade) that implement different DE variants, catering to a wide range of optimization needs. This flexibility allows users to select the most suitable DE strategy for their specific problem and dataset, ultimately leading to more accurate and reliable models.

dplbnDE has the potential to impact various domains where BNs are employed, such as healthcare, finance, and supply chain management, by improving the predictive capabilities and decision-making processes in these fields. Its potential for commercial applications, coupled with its versatility, paves the way for further adoption and innovation in data-driven industries.

As the dplbnDE package gains recognition and its user base expands, future developments could include incorporating additional optimization techniques, improving computational efficiency, and extending its applicability to other types of graphical models or machine learning algorithms. By continuing to evolve and address the needs of the research community and industry, dplbnDE has the potential to make a lasting impact on the study and application of BNs and optimization techniques.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data availability

No data was used for the research described in the article

Acknowledgments

A. Platas-López acknowledges support from the Mexican National Council of Science and Technology (CONACyT) and the Universidad Veracruzana, Mexico to pursue graduate studies at the Instituto de Investigaciones en Inteligencia Artificial.

References

- [1] Su C, Andrew A, Karagas MR, Borsuk ME. Using Bayesian networks to discover relations between genes, environment, and disease. *BioData Min* 2013;6(1). <http://dx.doi.org/10.1186/1756-0381-6-6>.
- [2] Ludwig M, Dührkop K, Böcker S. Bayesian networks for mass spectrometric metabolite identification via molecular fingerprints. *Bioinformatics* 2018;34(13):i333–40. <http://dx.doi.org/10.1093/bioinformatics/bty245>.
- [3] Fenton N, Neil M. Risk assessment and decision analysis with Bayesian networks. 2nd ed.. London, England: Chapman and Hall; 2018.
- [4] Anderson B. Using Bayesian networks to perform reject inference. *Expert Syst Appl* 2019;137:349–56. <http://dx.doi.org/10.1016/j.eswa.2019.07.011>.
- [5] Kaikkonen L, Parviainen T, Rahikainen M, Uusitalo L, Lehtikoinen A. Bayesian networks in environmental risk assessment: A review. *Integr Environ Assess Manag* 2020;17(1):62–78. <http://dx.doi.org/10.1002/ieam.4332>.

- [6] Pérez S, German-Labaume C, Mathiot S, Goix S, Chamaret P. Using Bayesian networks for environmental health risk assessment. *Environ Res* 2022;204:112059. <http://dx.doi.org/10.1016/j.envres.2021.112059>.
- [7] Platas-López A, Mezura-Montes E, Cruz-Ramírez N, Guerra-Hernández A. Discriminative learning of Bayesian network parameters by differential evolution. *Appl Math Model* 2021;93:244–56. <http://dx.doi.org/10.1016/j.apm.2020.12.026>.
- [8] Pearl J. Probabilistic reasoning in intelligent systems. Elsevier; 1988. <http://dx.doi.org/10.1016/c2009-0-27609-4>.
- [9] Friedman N, Geiger D, Goldszmidt M. Bayesian network classifiers. *Mach Learn* 1997;29:131–63.
- [10] Price K, Storn R. Minimizing the real functions of the ICEC'96 contest by differential evolution. In: *Proc. of IEEE C. evol. computat.*. 1996, p. 842–4.
- [11] Zhang J, Sanderson A. JADE: Adaptive differential evolution with optional external archive. *IEEE Trans Evol Comput* 2009;13:945–58.
- [12] Tanabe R, Fukunaga A. Improving the search performance of SHADE using linear population size reduction. In: *Proc. of IEEE C. evol. computat.*. 2014, p. 1658–65.
- [13] Scutari M. Learning Bayesian networks with the bnlearn R package. *J Stat Softw* 2010;35(3):1–22. <http://dx.doi.org/10.18637/jss.v035.i03>. URL <https://www.jstatsoft.org/index.php/jss/article/view/v035i03>.
- [14] Boettcher SG, Dethlefsen C. deal: A package for learning Bayesian networks. *J Stat Softw* 2003;8(20):1–40. <http://dx.doi.org/10.18637/jss.v008.i20>. URL <https://www.jstatsoft.org/index.php/jss/article/view/v008i20>.
- [15] Mullen KM, Ardia D, Gil DL, Windover D, Cline J. DEoptim: An R package for global optimization by differential evolution. *J Stat Softw* 2011;40(6):1–26. <http://dx.doi.org/10.18637/jss.v040.i06>. URL <https://www.jstatsoft.org/index.php/jss/article/view/v040i06>.
- [16] Kalisch M, Mächler M, Colombo D, Maathuis MH, Bühlmann P. Causal inference using graphical models with the R package pcalg. *J Stat Softw* 2012;47(11):1–26. <http://dx.doi.org/10.18637/jss.v047.i11>.
- [17] McLachlan S, Dube K, Hitman GA, Fenton NE, Kyrimi E. Bayesian networks in healthcare: Distribution by medical condition. *Artif Intell Med* 2020;107:101912. <http://dx.doi.org/10.1016/j.artmed.2020.101912>.
- [18] Kyrimi E, McLachlan S, Dube K, Neves MR, Fahmi A, Fenton N. A comprehensive scoping review of Bayesian networks in healthcare: Past, present and future. *Artif Intell Med* 2021;117:102108. <http://dx.doi.org/10.1016/j.artmed.2021.102108>.
- [19] Masmoudi K, Abid L, Masmoudi A. Credit risk modeling using Bayesian network with a latent variable. *Expert Syst Appl* 2019;127:157–66. <http://dx.doi.org/10.1016/j.eswa.2019.03.014>.
- [20] Hosseini S. A decision support system based on machine learned Bayesian network for predicting successful direct sales marketing. *J Manag Anal* 2021;8(2):295–315. <http://dx.doi.org/10.1080/23270012.2021.1897956>.
- [21] Gemela J. Financial analysis using Bayesian networks. *Appl Stoch Models Bus Ind* 2001;17(1):57–67. <http://dx.doi.org/10.1002/asmb.422>.
- [22] Neil M, Fenton N, Tailor M. Using Bayesian networks to model expected and unexpected operational losses. *Risk Anal* 2005;25(4):963–72. <http://dx.doi.org/10.1111/j.1539-6924.2005.00641.x>.
- [23] Chan LSH, Chu AMY, So MKP. A moving-window Bayesian network model for assessing systemic risk in financial markets. In: Zhao J, editor. *PLoS One* 2023;18(1):e0279888. <http://dx.doi.org/10.1371/journal.pone.0279888>.
- [24] Hosseini S, Ivanov D. Bayesian networks for supply chain risk, resilience and ripple effect analysis: A literature review. *Expert Syst Appl* 2020;161:113649. <http://dx.doi.org/10.1016/j.eswa.2020.113649>.
- [25] Villa S, Stella F. Bayesian networks for portfolio analysis and optimization. In: *Financial decision making using computational intelligence*. Boston, MA: Springer US; 2012, p. 209–32. http://dx.doi.org/10.1007/978-1-4614-3773-4_8.
- [26] Marrone S. Using Bayesian networks for highly available cloud-based web applications. *J Reliab Intell Environ* 2015;1(2–4):87–100. <http://dx.doi.org/10.1007/s40860-015-0009-z>.

CAPÍTULO 5

CONCLUSIÓN

El notable auge del modelado basado en agentes (MBA) y sus variantes evidencia la diversidad y riqueza de aplicaciones en múltiples disciplinas, desde la ecología hasta las ciencias sociales y la economía. Sin embargo, junto con este crecimiento vertiginoso, se ha presentado una notable fragmentación en términos de terminología, metodologías y protocolos de documentación. Esta diversidad terminológica, aunque refleja el eclecticismo del campo, puede generar confusión y redundancias, obstaculizando la consolidación y madurez del MBA como una herramienta robusta y ampliamente aceptada en la comunidad científica.

A medida que nuevos campos adoptan y adaptan el enfoque del MBA, se torna imperativo un esfuerzo colectivo para estandarizar terminologías, procesos y protocolos de documentación. La existencia de múltiples estándares, como TRACE, ODD, STRESS, y RAT-RS, sugiere que hay una necesidad palpable de una estructura común, pero también refleja la descentralización y el desarrollo paralelo en diferentes disciplinas. La unificación, o al menos una convergencia hacia prácticas compartidas, permitiría una comunicación más fluida entre investigadores, una comparación más sencilla entre estudios y, lo más importante, impulsaría el avance cohesivo del MBA.

Mientras que la diversidad en el modelado basado en agentes es un testimonio de

su versatilidad y aplicabilidad, el desafío pendiente radica en equilibrar esta riqueza con la claridad y coherencia. Es esencial que la comunidad de MBA trabaje colectivamente para forjar estándares y protocolos más unificados. Solo entonces podremos aprovechar al máximo el potencial del MBA, permitiendo que investigadores de diversas disciplinas colaboren con eficacia, compartan sus hallazgos y contribuyan a un cuerpo de conocimiento más consolidado y enriquecido.

La naturaleza intrínseca de los Modelos Basados en Agentes, si bien ofrece una versatilidad sin precedentes en la representación de sistemas complejos, no está exenta de críticas y limitaciones. En este trabajo, reconocimos estas limitaciones y buscamos superarlas con el auxilio de técnicas avanzadas de aprendizaje automático supervisado.

La incorporación de técnicas de aprendizaje supervisado en la modelación y simulación basadas en agentes emergió no sólo como una solución técnica, sino también como una estrategia para repensar y refinar el proceso de diseño, calibración y análisis de estos modelos. Las críticas tradicionales relacionadas con la complejidad en diseñar un MBA, su difícil calibración/validación, y el desafío de analizar/interpretar los resultados, en gran medida, pueden ser mitigadas al emplear dichas técnicas. Esta premisa, asumida como hipótesis en los albores de este trabajo, ha guiado el desarrollo del mismo.

Con base en la hipótesis planteada y utilizando una metodología definida, abordamos la tarea de alcanzar los objetivos establecidos al inicio de la investigación. Las conclusiones presentadas a continuación derivan directamente de estos objetivos, demostrando la consistencia y rigurosidad del estudio realizado.

Revisión Crítica de la Literatura (O.E. 1):

Se realizó una revisión exhaustiva de la literatura, identificando una variedad de aplicaciones de ML en ABM. Esta revisión reveló una clara distinción entre los enfoques en línea y fuera de línea del ML, proporcionando una base sólida para una clasificación estructurada y coherente de la literatura existente.

Propuesta de un Marco de Referencia (O.E. 2):

Basado en la revisión de la literatura, se propuso un marco de referencia que integra técnicas de aprendizaje supervisado en ABM. La extensión del Protocolo ODD, en particular, llenó un vacío en la documentación de ABM al incorporar aspectos relacionados con el aprendizaje, favoreciendo así una mayor replicabilidad y comprensión en futuras investigaciones.

Implementación de un Caso de Estudio (O.E. 3):

A través de un caso de estudio sobre la evasión del impuesto sobre nóminas en México, se demostró la eficacia de integrar un modelo de ML en una simulación ABM. Esta implementación reveló dinámicas y factores que habrían sido difíciles de identificar mediante enfoques analíticos tradicionales.

Aplicación de Técnicas de Caja Blanca (O.E. 4):

Las técnicas de caja blanca, en particular las redes bayesianas, fueron empleadas para ofrecer una interpretación transparente de los resultados del modelo. Esta aproximación no solo fortaleció la validación y análisis de los resultados, sino que también subrayó el potencial de tales técnicas para abordar problemas complejos en ABM.

A pesar de todos estos alcances, es esencial reconocer que la integración de aprendizaje supervisado en MBA no es una solución mágica. La decisión de optar por modelos de caja negra o caja blanca es una elección meticulosa, que exige equilibrar precisión con transparencia, siempre guiada por el objetivo subyacente de la investigación.

En la actualidad, se observa una tendencia creciente hacia el uso de redes neuronales en el ámbito del MBA, una inclinación que, en algunas situaciones, podría no ser la más apropiada. Es posible que este fenómeno se deba a la popularidad y al auge de las redes neuronales más que a una elección informada basada en la adecuación técnica. Dado el carácter intrínsecamente interdisciplinario del MBA, la elección del modelo ML podría estar siendo influenciada más por tendencias actuales que por una comprensión profunda de las necesidades y objetivos de la tarea específica.

Para contrarrestar este fenómeno, proponemos varias estrategias. Primero, es esencial

fortalecer las comunidades interdisciplinarias, creando espacios donde los expertos en MBA y ML puedan colaborar y compartir conocimientos. Esta colaboración fomentaría la formación de equipos que combinen lo mejor de ambos mundos, permitiendo decisiones más informadas sobre qué herramientas utilizar en función de la naturaleza del problema en cuestión.

Adicionalmente, la educación y la formación continua juegan un papel crucial. Al familiarizar a los profesionales del MBA con los fundamentos y limitaciones del ML, y viceversa, se facilitaría una integración más efectiva y coherente de estas disciplinas. Las decisiones dejarían de estar basadas en modas y estarían más arraigadas en la comprensión técnica y teórica.

También, es vital que la comunidad científica adopte un enfoque más crítico y reflexivo, cuestionando continuamente las herramientas y métodos que se utilizan, para garantizar que las elecciones hechas sean las más adecuadas para el problema en cuestión y no simplemente una adopción pasiva de las últimas tendencias.

Además, aunque el aprendizaje supervisado ofrece mejoras en el diseño, calibración y análisis de MBA, no aborda todas las críticas hacia estos modelos. Por ejemplo, la necesidad de habilidades de programación sigue siendo una barrera para muchos investigadores, además, el coste computacional podría ser un tema que impida ciertas implementaciones.

Es innegable que tanto el costo computacional como las habilidades de programación representan barreras latentes en la implementación de MBA. El costo computacional, influenciado por el teorema del “no free lunch”, se manifiesta en muchas técnicas avanzadas, no solo en MBA. Esta dicotomía entre la capacidad de representar fenómenos complejos y el costo computacional asociado, nos lleva a reflexionar sobre la esencia de los modelos.

Herramientas como los modelos DSGE enfrentan desafíos similares, lo que sugiere que una reflexión más profunda sobre la esencia y complejidad de los modelos es necesaria. En la búsqueda constante de realismo, es posible que hayamos perdido de vista la simplicidad y eficiencia que un modelo puede ofrecer. Es esencial que se fomente una mentalidad crítica entre los científicos: ¿es imperativo un modelo tan detallado? Adicionalmente, podríamos explorar técnicas que prioricen la eficiencia computacional desde la concepción del modelo, así como la implementación de heurísticas o algoritmos optimizados.

En cuanto a la programación, la realidad actual puede parecer desalentadora para muchos.



Sin embargo, la solución podría estar en el desarrollo de herramientas y plataformas más intuitivas. Imaginar soluciones de software que simplifiquen la creación de modelos, como interfaces de “drag and drop” o lenguajes más amigables, podría democratizar el uso de MBA. A pesar de esto, es crucial subrayar la importancia de que los usuarios tengan una comprensión sólida de las metodologías y técnicas que están utilizando. La facilidad en la implementación nunca debe eclipsar la necesidad de un entendimiento profundo y justificado.

Finalmente, exhortamos a la comunidad a adoptar, adaptar y profundizar en el flujo de trabajo propuesto, reconociendo el vasto potencial que la fusión del aprendizaje supervisado y MBA ofrece para la innovación en la modelación de sistemas complejos. A medida que el campo del ABM continúa creciendo, la implementación de nuestra extensión sugerida al protocolo ODD en variados dominios y contextos se destaca como una veta prometedora para futuras investigaciones. Paralelamente, se invita a una revisión más profunda que destaque tendencias emergentes y fronteras aún por descubrir en la modelización basada en agentes fortalecida por el aprendizaje automático.

ANEXO

La evasión del impuesto sobre la nómina en México se modela como un estudio de caso para ejemplificar los beneficios de la combinación de aprendizaje automático y MBA. Cada estado mexicano tiene autonomía sobre la forma en que se recauda el impuesto sobre la nómina. Por lo tanto, para modelar estos diferentes escenarios fiscales y sus efectos, se realiza una representación explícita del espacio a través de un Sistema de Información Geográfica con teselación hexagonal.

A.1. El modelo de evasión de impuestos sobre la nómina

En aras de la reproducibilidad, los detalles del modelo se describirán siguiendo el protocolo descrito por [Grimm et al. \(2010\)](#) denominado ODD (Overview, Design Concepts and Details), que está organizado en tres partes:

1. Visión general. Una descripción general del modelo, incluyendo su propósito y sus componentes básicos: agentes, variables que los describen y el entorno, además de las escalas utilizadas en el modelo, por ejemplo, tiempo y espacio; así como una descripción general de los procesos y su programación.

2. Conceptos de diseño. Una breve descripción de los principios básicos que subyacen al diseño del modelo, por ejemplo, racionalidad, emergencia, adaptación, aprendizaje, etc.
3. Detalles. Definiciones completas de los submodelos involucrados.

A.1.1. Visión general

Propósito. Analizar el efecto de los determinantes institucionales sobre la evasión del impuesto sobre la nómina en México.

Entidades, variables de estado y escalas.

- Agentes: Empleadores y autoridad fiscal.
- Medio ambiente: una representación a nivel estatal mexicano del sistema de impuestos sobre la nómina.
- Escalas: El tiempo se representa en períodos discretos, representando cada paso un mes, que corresponde al período de recaudación de impuestos según la legislación vigente. Cada agente empleador en el modelo representa 2,000 empleadores en el mercado laboral mexicano de 2019, se eligió esta escala para tener un margen de error menor al 5 % con un intervalo de confianza del 99 % en la muestra seleccionada,
- Variables de estado: Los atributos que caracterizan a cada agente se muestran en la Tabla A.1.

Agente	Tipo	Atributos
Auditor	Agset	employers-of-auditor
	Float	penalty-collected; tax-collected
	Int	ent-auditor
Employer	Bool	audit?; audited?; formal-or-informal (mh_col)
	Float	corruption; declared-tax; insecurity; payroll; payroll*; prob-formal; production; risk-aversion- ρ ; tax; undeclared-payroll; undeclared-tax; α_s ; δ
	Int	eda (age); business-size (ambito2); education (anios_esc); economic-activity (c_ocu11c); state (ent); income (ing7c); size-of-region (t_loc); type-of-taxpayer

Tabla A.1: Variables de estado por tipo de agente. Entre paréntesis, el nombre del atributo en el conjunto de datos del INEGI.

Resumen y programación del proceso.

1. Los empleadores deciden en qué mercado producir, formal o informal, consultando un modelo de aprendizaje automático.
2. Los empleadores informales son evasores totales. Los empleadores formales calculan la cantidad de impuestos a declarar. Si calculan declarar impuestos cero, también se convierten en evasores totales. Si calculan declarar todos los impuestos, se convierten en contribuyentes completos. En cualquier otro caso, se convierten en evasores parciales.
3. La autoridad fiscal recauda la cantidad declarada de impuestos.
4. La autoridad fiscal realiza auditorías de forma aleatoria. Si la auditoría tiene éxito, los evasores parciales y totales deben pagar la cantidad evadida y una multa por la cantidad no declarada.
5. Cada 12 meses los empleadores aumentan su edad. Con cierta probabilidad, en cada período, los empleadores pueden morir. Si esto sucede, son reemplazados por otro empleador con las mismas características excepto por la edad.

A.1.2. Conceptos de diseño

Principios básicos.

Los ingresos declarados voluntariamente, aumentan con el aumento de los ingresos individuales, la tasa de penalización y la probabilidad de auditoría, y disminuyen con el aumento de la tasa de impuestos, como se propone en la teoría neoclásica de la evasión de [Allingham y Sandmo \(1972\)](#), y también incluye Suministro exponencial de servicios públicos y bienes públicos descrito por [Hokamp \(2014\)](#).

La distribución de la producción entre los empresarios sigue una ley de potencias en la que existe una marcada desigualdad en la distribución, caracterizada por un pequeño porcentaje de personas que poseen la mayoría de los recursos ([Chakraborti y Patriarca, 2008](#)), lo que caracteriza a una economía capitalista como la de México. La mortalidad de los empleadores sigue una distribución de Weibull ([Pinder et al., 1978](#)). Esta función se ajusta para el caso de México, donde la esperanza de vida al nacer es de setenta y cinco años ([CONAPO, 2020](#)).

Emergencia.

La extensión de la evasión fiscal es un resultado emergente del comportamiento adaptativo de los empleadores. Se espera que estos resultados varíen de manera compleja cuando cambien las características de los individuos o su entorno. Otros resultados, como la distribución de la producción y la edad, son impuestos más estrictamente por reglas y, por lo tanto, dependen menos de lo que hacen los individuos. Sin embargo, estos resultados son importantes para la validación del modelo.

Adaptación.

Los empleadores pueden cambiar su alineación en el sector formal. Para tomar esa decisión, consultan un modelo de aprendizaje automático previamente entrenado. Con este rasgo, los agentes podrían convertirse en evasores totales y, por lo tanto, evitar pagar impuestos. La autoridad tributaria no cuenta con mecanismos de adaptación.

Objetivos.

El objetivo del empleador es aumentar su utilidad pagando la menor cantidad de impuestos. El objetivo de la autoridad fiscal es aumentar la recaudación de impuestos.

Aprendizaje.

Con la propuesta de extensión para la inclusión del aprendizaje se responde a las siguientes interrogantes. *¿Cuál es el propósito del aprendizaje?* El propósito del aprendizaje es inducir las reglas que definen el comportamiento de los agentes a partir de los datos existentes. *¿Cuándo se realiza el aprendizaje?* El aprendizaje se realiza antes de la simulación para obtener modelos para generar las reglas de los agentes incluidos en el ABM. *¿Qué componentes de la ABM se ven afectados por el aprendizaje?* El entorno se ve afectado al encapsular en un modelo aprendido los datos del sistema fiscal y la calidad de los bienes públicos. *¿Cómo se calcula el aprendizaje?* El aprendizaje se calcula mediante un algoritmo supervisado.

Predicción.

El aprendizaje de los empleadores es un proceso fuera de línea. Durante la simulación, los empleadores perciben sus condiciones actuales y actúan en consecuencia. Los empresarios no cuentan con un modelo interno para estimar sus condiciones futuras o las consecuencias de su decisión.

Sensado.

Los empleadores consideran el tamaño de su empresa, educación, sector de ocupación, estado, tamaño de la localidad, edad e ingresos como estados internos, y perciben la tasa impositiva y el nivel de inseguridad como variables ambientales. Si una auditoría tiene éxito, la autoridad fiscal puede cobrar impuestos no declarados a los empleadores.

Interacción.

Las interacciones entre los empleadores y la autoridad tributaria son directas, cuando con cierta probabilidad se lleva a cabo una auditoría.

Estocasticidad.

El proceso de auditoría se asume aleatorio, tanto la probabilidad de llevarlo a cabo como su probabilidad de éxito siguen una distribución uniforme. El proceso de inicialización de los ingresos de los empleadores también se considera aleatorio, siguiendo una ley de potencia. La probabilidad de muerte de los empleadores también es aleatoria siguiendo una distribución de supervivencia de Weibull. En el primer proceso, la estocasticidad se usa para hacer que los eventos ocurran con una frecuencia específica. En los dos últimos procesos, la estocasticidad se utiliza para reproducir la variabilidad en procesos para los que no es importante modelar las causas reales de la variabilidad.

Colectivos.

Los empleadores se agrupan en tres tipos de comportamiento tributario: evasores totales, evasores parciales y contribuyentes completos. Este colectivo es una definición del modelo en el que el conjunto de empleadores con ciertas propiedades sobre la cantidad de impuestos pagados se define como un tipo de empleador separado con sus propias variables y rasgos.

Observación.

Al final de cada ejecución, se recopilan datos sobre el alcance de la evasión fiscal, la cantidad de impuestos recaudados y el número de empleadores evasores totales.

A.1.3. Detalles**Inicialización.**

Inicialización. En el momento $t = 0$, de cada ejecución de simulación, se inicializan $N = 1337$ empleadores y se distribuyen en cada estado de acuerdo con los datos de entrada. También se inicializan 32 auditores en representación de cada autoridad fiscal estatal. Algunas variables de estado se inicializan mediante un submodelo, ya sea determinista o aleatorio, mientras que otras se basan en los datos que se muestran en la tabla A.2.

Tabla A.2: Variables de estado y método de inicialización. Entre paréntesis, está el nombre del atributo en el conjunto de datos del Instituto Nacional de Estadística y Geografía (INEGI) (INEGI, 2019a) de variables de estado.

Agent	Attributes	Type	Initialization	Value
Auditor	penalty-collected	Float	Deterministic	0
	tax-collected	Float		0
	my-employers	AgSet	Random	Submodel 5
	ent-auditor	Int		[1, 32]
Employer	business-size (ambito2)	Int	Database	{0, 2, 3, 4, 5, 8}
	education (anios_esc)	Int		[0, 20]
	economic-activity (c_ocu11c)	Int		[1, 10]
	age (eda)	Int		[17, 98]
	mexican-state (ent)	Int		[1, 32]
	income (ing7c)	Int		[1, 7]
	formal-or-informal (mh_col)	Int		[0, 1]
	size-of-region (t_loc)	Int	Deterministic	[1, 4]
	corruption	Float		(0, 1)
	insecurity	Float		(0, 1)
	tax	Float		(0, 1)
	audit?	Bool		false
	audited?	Bool		false
	type-of-taxpayer	Int		2
	declared-tax	Float		0
	payroll	Float		Submodel 11
	payroll*	Float		Submodel 12
	risk-aversion- ρ	Float		Submodel 16
	undeclared-payroll	Float		0
	undeclared-tax	Float		0
	α -s	Float		0.05
	δ	Float		-0.1
	prob-formal	Float	Random	(0, 1)
	production	Float		Submodel 10

De la misma manera, los parámetros de entrada de la simulación se adoptaron de la literatura, base de datos o mediante experimentación. Table A.3 muestra los valores iniciales del modelo de referencia. En México, la tasa de penalización varía en algunos estados entre 55 y 75 %, y en otros entre 75 y 100 %, por lo que se adoptó un valor inicial de 75 %. La variación de tasas impositivas, percepción de inseguridad y corrupción parten de cero, es

decir se toma el valor señalado en la legislación tributaria local y en el INEGI. De acuerdo con [Bonet y Rueda \(2012\)](#), la efectividad de la recaudación de impuestos en México es de aproximadamente 0.7. Para el umbral de decisión se utilizó 0.5, ya que según ([Witten et al., 2017](#)), el umbral de 0.5 se usa comúnmente para estimar la probabilidad de un conjunto con dos clases. Para la probabilidad de auditoría y la efectividad de las auditorías se encontraron valores a través de la experimentación para garantizar salidas similares a la dinámica del sistema tributario mexicano, el cual es pequeño pero eficiente. La interfaz de usuario también tiene otros parámetros de entrada, uno de ellos permite active y desactive el modelo de aprendizaje automático.

Tabla A.3: Input parameter initialization of baseline model.

Parameter	Description	Value	Initialization
π	Penalty rate	0.75	Database
α	Audit probability	0.05	Experimentation
ϵ_{AP}	Effectiveness of audit process	0.75	Experimentation
ϵ_{TC}	Effectiveness of tax collection	0.70	Literature Bonet y Rueda (2012)
$\Delta\theta$	Variation in tax rate	0.00	Database
ΔPI	Variation in perceived insecurity	0.00	Database
ΔPC	Variation in perceived corruption	0.00	Database
τ	Threshold for formal or informal sector choice	0.50	Literature Witten et al. (2017)

Parámetro	Descripción	Valor	Inicialización
π	Tasa de penalización	0.75	Datos
α	Probabilidad de auditoría	0.05	Experimental
ϵ_{AP}	Efectividad del proceso de auditoría	0.75	Experimental
ϵ_{TC}	Efectividad de la recaudación fiscal	0.7	Bonet y Rueda
$\Delta\theta$	Variación en tasa impositiva	0.0	Datos
ΔPI	Variación en inseguridad percibida	0.0	Datos
ΔPC	Variación en corrupción percibida	0.0	Datos
τ	Umbral para elección formal-informal	0.5	Witten et al.

Tabla A.4: Inicialización del parámetro de entrada del modelo de línea base.

Datos de entrada.

¿El modelo utiliza datos de fuentes externas (archivos de datos u otros modelos) para

representar algún elemento del modelo? Sí, el modelo utiliza datos de entrada externos para conocer el sistema tributario mexicano, es decir, la Encuesta Nacional de Ocupación y Empleo (ENOE), la Encuesta Nacional de Calidad e Impacto Gubernamental (ENCIG), así como las leyes tributarias de los distintos estados donde se especifica la tasa del impuesto sobre la nómina. El periodo considerado para las 3 fuentes de información es del 2011 al 2019.

La ENOE es una encuesta trimestral, pero solo se utiliza el tercer trimestre como referencia para la información anual. Entre otros datos, podemos conocer distintas variables sociodemográficas sobre las características de los empresarios. A partir de estas variables, se puede determinar si un empleador está en el sector formal o informal. ENCIG es una encuesta semestral en la que se pregunta a las personas sobre los 3 (tres) principales problemas que creen que existen en su estado. Los principales problemas entre todos los estados son la inseguridad, la corrupción y el desempleo. La inseguridad y la corrupción se resumen por estado y se interpolan para obtener información anual. Los datos resultantes se unen al conjunto de datos ENOE junto con los datos fiscales. El conjunto de datos recolectado de esta manera da una matriz de tamaño 71833×10 , donde la proporción de empleadores formales es 60.57 %.

¿Se modelan los datos con aprendizaje automático? Sí, a partir de los datos ENOE preprocesados resultantes, el muestreo se realiza mediante el método de pivote local con el algoritmo ofrecido por el paquete R “SamplingBigData” (Lisic y Cruze, 2016), que genera efectivamente un conjunto de datos de muestra equilibrado. Los atributos seleccionados para generar la muestra balanceada fueron el estado, la clasificación del empleador por pertenecer al sector formal o informal y el tamaño de la unidad económica del empleador. Esto asegura que los empleadores en el modelo reflejen las proporciones reales del mercado laboral mexicano. El tamaño de la muestra seleccionada corresponde a una escala de 1 a 2000. La figura A.1 muestra el proceso de integración de ML.

Submodelos

1. Luego de realizar un pre-procesamiento de la base de datos incluyendo una selección manual y una eliminación recursiva de características con el algoritmo de remuestreo (Gościk y Łukaszuk, 2013) disponible en el paquete R “caret” (Kuhn, 2008), se determinó que las principales variables que determinan el sector del empleador son el tamaño de

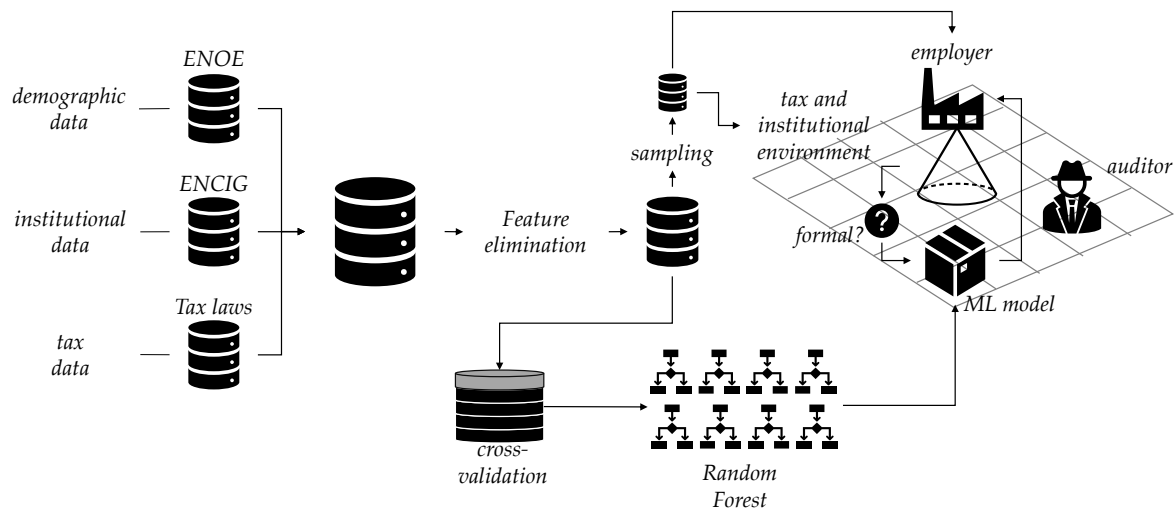


Figura. A.1: Diagrama de integración de ML en el Modelo de Evasión de Impuestos sobre Nómina.

- la empresa (ambito2), educación (anios_esc), actividad económica (c_ocu11c), estado (ent), tamaño de la región (t_loc), y edad (eda).
- Se eligió una implementación rápida de Random Forest en el paquete R “ranger”(Wright y Ziegler, 2017) para aprender de los datos porque proporciona un ajuste y evaluación rápidos del modelo, es resistente a los valores atípicos, puede manejar asociaciones lineales simples y no lineales complicadas , y produce una precisión de predicción competitiva (Han y Kim, 2021). Para ajustar los hiperparámetros y evaluar el desempeño del modelo, se realizó una validación cruzada con $k = 10$ pliegues. Los valores finales de los hiperparámetros fueron $mtry = 20$, $ntrees = 100$ y $nodesize = 1$. Ese ajuste consigue una precisión del 83,79 %, lo que se considera bueno para evitar el sobreajuste. El modelo entrenado estará disponible para los empleadores durante la simulación.
 - Se carga una capa del Sistema de Información Geográfica. Cada polígono es un teselado hexagonal del estado mexicano correspondiente.
 - Se generan $N = 1337$ empleadores y se inician con información de la base de datos y se trasladan a su estado correspondiente.
 - Los auditores se generan y ubican en su estado asignado.
 - Se carga el modelo de bosque aleatorio aprendido a priori.
 - Genere valores de la ley de Pareto con la siguiente función de distribución.

$$f(x) \sim x^{-1-\gamma}$$

8. Donde γ se conoce como el exponente de Pareto y se estima que es $\approx 3/2$ para caracterizar una economía capitalista ([Chakraborti y Patriarca, 2008](#)).
9. x son los valores generados por una función de distribución normal con media 2 y desviación estándar de 0.2 para los empleadores informales y 0.3 para los formales.
10. Asignar un valor fijo de producción mensual a cada empleador. Los valores de ley de potencia generados se multiplican por 23 en el caso de los empleadores informales y por 50 para los formales. Esas cantidades generan una distribución de Pareto perfectamente mixta de acuerdo a los principios básicos y preservan la participación de la economía informal en el Producto Interno Bruto (PIB) mexicano ([INEGI, 2019c](#)).
11. Para simplificar, se supone que cada empleador asigna el 30 por ciento del valor de la producción a la nómina W . La participación de los salarios en el PIB mexicano está entre el 30 y el 40 % ([Samaniego Breach, 2014](#)).
12. Al comienzo de la simulación, se supone que los empleadores no informales declaran todo el impuesto, es decir, declaran la nómina $W^* = W$.
13. Al principio, el impuesto declarado X^* por cada empleador es igual a la nómina declarada W^* multiplicada por la tasa impositiva θ en su estado.
14. Cada 12 periodos (meses) los empleadores aumentan su edad, y consultan el modelo aprendido para decidir si optan por el mercado formal o informal, tomando como referencia sus atributos internos y la inseguridad percibida.
15. Los empleadores informales no declaran impuestos.
16. Por norma social ([Hokamp, 2014](#)), los empleadores modifican su aversión al riesgo ρ según su edad de la siguiente manera:

$$\rho \sim \begin{cases} U(0,0,0,25) & \text{if } \text{age} \leq 34 \\ U(0,25,0,5) & \text{if } 34 < \text{age} \leq 51 \\ U(0,5,0,75) & \text{if } 51 < \text{age} \leq 67 \\ U(0,75,1,0) & \text{if } \text{age} \geq 67 \end{cases}$$

17. Sea β la eficiencia percibida de los bienes públicos y π la tasa de penalización.
18. Sea ϵ_{AP} y ϵ_{TC} la efectividad del proceso de auditoría y la recaudación de impuestos respectivamente.
19. Sea α la verdadera probabilidad de auditoría y α_S la probabilidad subjetiva de auditoría conocida por el empleador.
20. Sea $\delta = 0,1$, el parámetro de actualización de α_S .
21. Si un empleador es auditado en un período específico, la probabilidad de auditoría subjetiva se convierte en 1.
22. En cada período (si no se vuelve a auditar) α_S disminuye en δ hasta $\alpha_S = \alpha$.
23. Cada período, los empleadores calculan la cantidad de impuestos a declarar voluntariamente X^* , aplicando el procedimiento de maximización de la utilidad esperada adoptado por [Allingham y Sandmo \(1972\)](#). Sea límite inferior:

$$\alpha_S > \frac{1}{1 + \left(\frac{(1-\beta(1-\epsilon_{AP}))\pi}{(1-\beta(1-\epsilon_{TC}))\theta} - 1 \right) e^{\rho(1-\beta(1-\epsilon_{AP}))(\pi W)}}$$

24. Y el límite superior:

$$\alpha_S < \frac{1}{1 + \left(\frac{(1-\beta(1-\epsilon_{AP}))\pi}{(1-\beta(1-\epsilon_{TC}))\theta} - 1 \right) e^{\rho(1-\beta(1-\epsilon_{AP}))(\pi W)}}$$

25. Si la probabilidad de auditoría subjetiva α_S excede el límite superior en el submodelo 22, el empleador se convierte en totalmente cumplidor de impuestos, es decir, $X^* = W\theta$, y cuando α_S cae por debajo del límite inferior en el submodelo 21, el empleador evade por completo, es decir $X^* = 0$.

26. Para α_S en el rango de una solución interna, el empleador declara voluntariamente:

$$X^* = W - \frac{\ln \left(\frac{(1-\alpha_S)(1-\beta(1-\epsilon_{TC}))\theta}{\alpha_S((1-\beta(1-\epsilon_{AP}))\pi - (1-\beta(1-\epsilon_{TC}))\theta)} \right)}{\rho\pi(1-\beta(1-\epsilon_{AP}))}$$

27. La autoridad fiscal recauda los impuestos sobre la nómina que los empleadores declararon voluntariamente.
28. La autoridad fiscal realiza auditorías con una probabilidad aleatoria de α y un nivel de efectividad de ϵ_{AP} .
29. Si se detecta un evasor, se recauda el impuesto no declarado y se aplica una tasa de penalización π sobre el impuesto no declarado.
30. En cada período, los empleadores tienen una probabilidad de morir, siguiendo una función de derivación de cuantiles de Weibull:

$$Q(p) = \lambda \left[\frac{1}{1-p} \right]$$

31. Donde $\lambda = 0,019$ y $k = 0,479$ son los parámetros de escala y forma respectivamente.
32. Se supone que cuando fallece un empleador, toma su lugar otra persona con los mismos atributos, excepto la edad, que se genera de acuerdo a:

$$eda = \lfloor X \rfloor$$

$$X \sim N(\mu, \sigma^2) \sim N(37, 6)$$

33. En cada tiempo t , la salida observada Extensión de la evasión fiscal (ETE) se calcula de la siguiente manera:

$$ETE_t = 1 - \frac{\sum_{i=1}^N W^*}{\sum_{i=1}^N W}$$

A.2. Resultados

El modelo fue implementado en Netlogo ([Wilensky, 1999](#)). La figura [A.2](#) muestra parte de la GUI resultante que permite la inicialización de parámetros globales y proporciona una vista de los agentes en un entorno de cuadrícula. El color de los hexágonos representa el impuesto recaudado por la entidad. Los agentes de empleo son de color azul, rojo o cian, dependiendo de si cumplen con los impuestos, evaden o evaden parcialmente, respectivamente. Los monitores del lado derecho muestran las salidas producidas en el modelo con la configuración de parámetros dada.

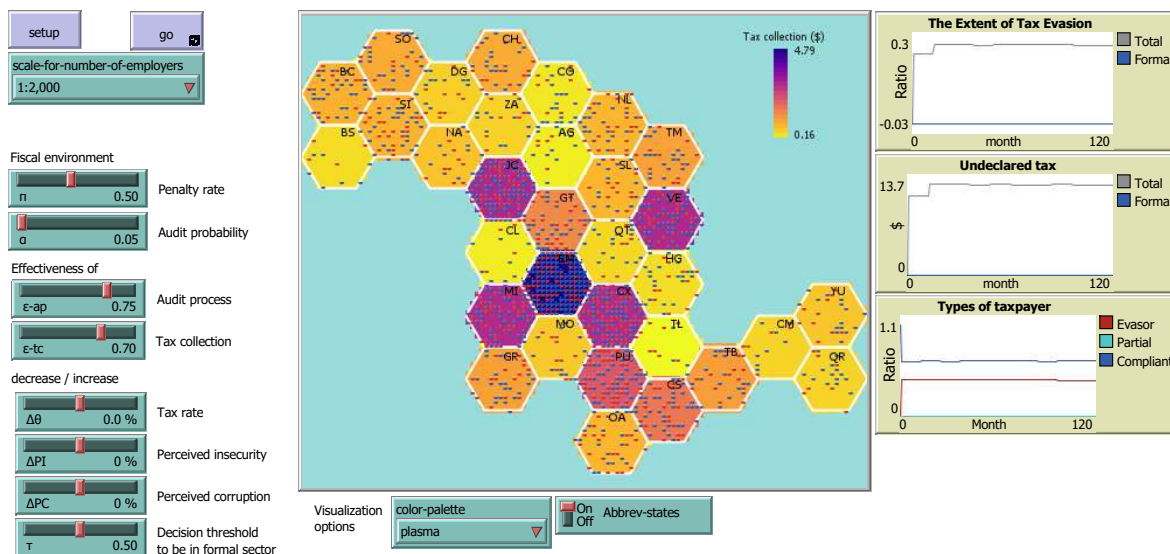


Figura. A.2: Interfaz gráfica de usuario del modelo..

Aunque un modelo de simulación demostrativo es útil para los estudios hipotéticos que exploran diferentes escenarios, el análisis de políticas requiere modelos de simulación descriptiva validados (Marks, 2013). En este apartado se realiza una validación de los resultados del modelo de línea base en términos micro y macro. La validación consiste en replicar las propiedades estadísticas de los datos observados. Después de la validación, se realiza una comparación entre las salidas de simulación de los modelos con y sin aprendizaje automático. La evaluación se realiza en términos del alcance de la evasión fiscal, los impuestos recaudados y la proporción de empleadores que evaden por completo.

A.2.1. Validación

La validación empírica de los modelos económicos basados en agentes ha experimentado un desarrollo significativo (Fagiolo et al., 2019). Entre las diferentes propuestas de validación, el enfoque indirecto es una de las más utilizadas y la primera evaluación que debe satisfacer un ABM. El núcleo de la validación indirecta es que las variables micro y macro pueden generar hechos estilizados o propiedades estadísticas que el modelador puede comparar con los obtenidos del análisis empírico del conjunto de datos del mundo real correspondiente (Windrum et al., 2007; Fagiolo et al., 2019). Se utilizará un enfoque indirecto para realizar la validación ya que existen variables de salida que deben seguir hechos estilizados así como datos reales para realizar la comparación.

A nivel micro, la validación consiste en verificar la existencia de algunos hechos estilizados en economía. Como se mencionó anteriormente, uno de los principios básicos radica en el supuesto de que en una economía capitalista la distribución de la producción sigue una ley de potencia. Y esta distribución también se traslada a la distribución de la nómina, asumiendo que una mayor producción demanda una mayor cantidad de salarios. La Figura A.3 muestra que los submodelos 5 a 9, que se adhieren al principio básico de la distribución de Pareto, generan una ley de potencia en la distribución de la nómina, es decir, un sesgo positivo en la distribución. En la Figura A.4, el sesgo positivo de 25 corridas independientes se mide con el *skewness*, una medida de la asimetría de la distribución de probabilidad de una variable aleatoria de valor real sobre su media (Joanes y Gill, 1998).

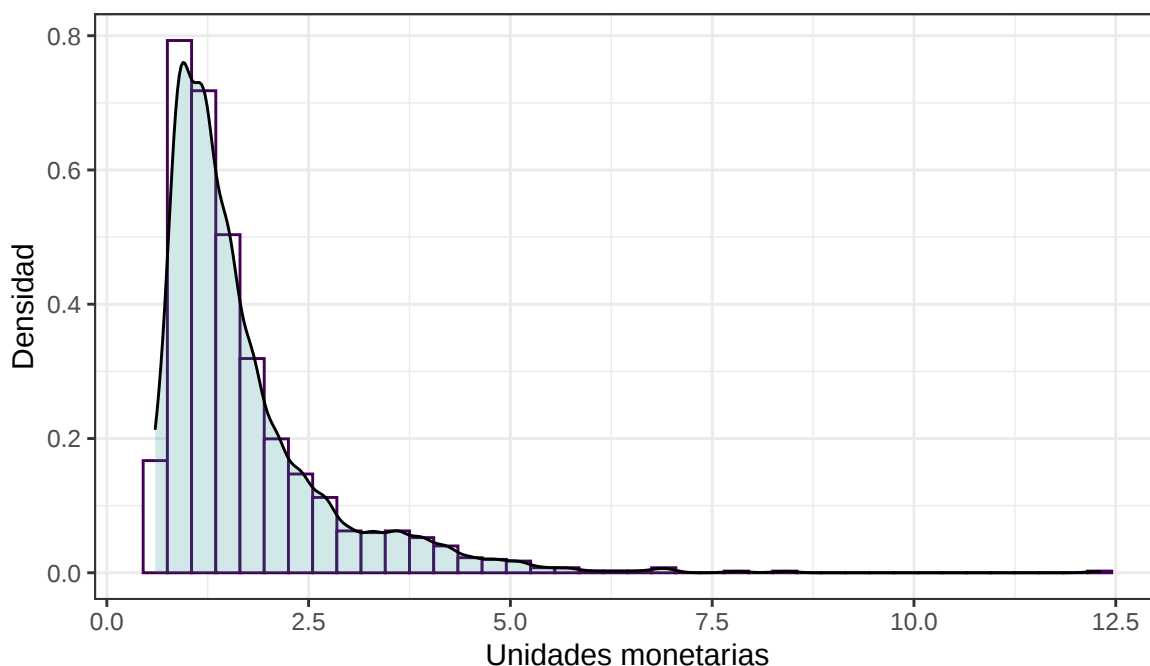


Figura. A.3: La distribución de la nómina sigue una ley de potencias.

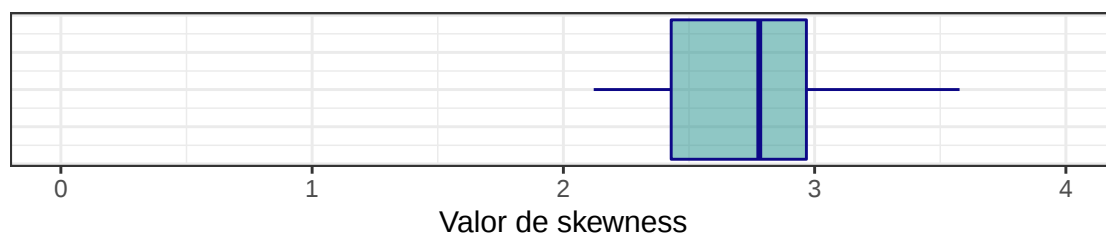


Figura. A.4: Valores de skewness obtenidos en 25 corridas independientes sobre la distribución de la nómina. Se considera que los valores superiores a 1 corresponden a distribuciones muy sesgadas positivamente.

A nivel macro, se realizará una comparación entre los montos de recaudación de impuestos obtenidos en la simulación y los reflejados en las estadísticas oficiales por estado. Los datos sobre recaudación de impuestos que tienen como fuente la simulación, se dividen en simulación con aprendizaje automático “on” y “off”. Para los datos de simulación, los valores corresponden a la media aritmética de 25 corridas independientes del último año simulado. Los valores de la fuente real corresponden a los reportados por [INEGI \(2019b\)](#).

Dado que la fuente de datos “real” contiene cantidades mucho mayores que las producidas en el modelo, todos los datos se centran y se escalan para que sean comparables. La Figura A.5 muestra la distribución de la recaudación de impuestos sobre la nómina por fuente de datos. Visualmente, la distribución de la recaudación del impuesto sobre la nómina se aproxima a los datos encontrados en las estadísticas.

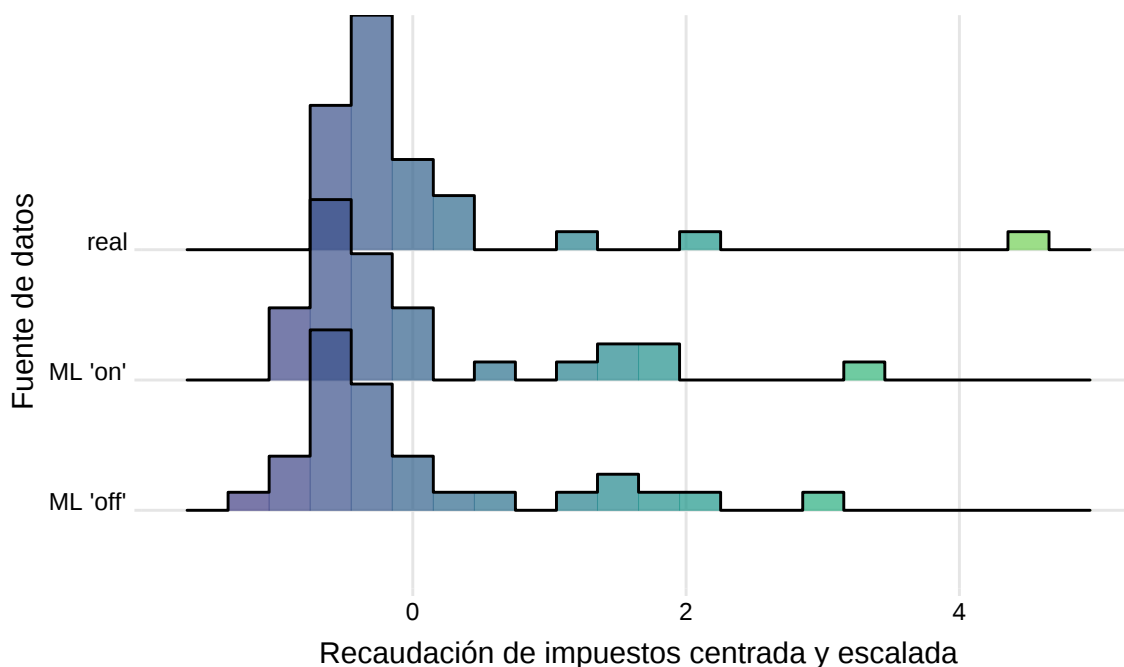


Figura. A.5: Distribución de la recaudación de impuestos sobre la nómina escalada y centrada por fuente de datos.

Los valores de la recaudación estatal del impuesto sobre la nómina de las estadísticas para el período 2010 a 2019 y de la media aritmética de 25 corridas independientes del último año simulado con y sin aprendizaje automático se clasifican de menor a mayor. La relación monótona entre estas fuentes de datos se mide con el coeficiente de correlación

de Spearman. En política, los valores absolutos mayores o iguales a 0.6 se consideran una fuerte correlación (Akoglu, 2018).

La Tabla A.5 muestra los coeficientes de correlación de Spearman entre estos datos. Se obtiene una fuerte correlación para todos los años considerados. Tanto la validación micro como la macro demuestran que el modelo de línea de base es una herramienta prometedora para analizar la evasión de impuestos sobre la nómina.

Modelo	Datos reales por año									
	2010	2011	2012	2013	2014	2015	2016	2017	2018	2019
ML 'off'	0.73	0.75	0.68	0.68	0.64	0.61	0.61	0.6	0.6	0.61
ML 'on'	0.71	0.73	0.68	0.68	0.65	0.62	0.6	0.6	0.6	0.61

Tabla A.5: Coeficiente de correlación de Spearman entre diferentes años de recaudación de impuestos sobre la nómina y los modelos con aprendizaje (ML 'on') y sin aprendizaje (ML 'off').

Finalmente, se calcula el error cuadrático medio (RMSE) para validar los resultados de los modelos con y sin aprendizaje automático en comparación con los valores reales. La figura A.6 resume por estado mexicano y modelo, el error entre la recaudación anual de impuestos prevista en 10 años simulados y la recaudación real en el período 2010 a 2019. El RMSE en el modelo con aprendizaje automático es menor en 19 de los 32 estados.

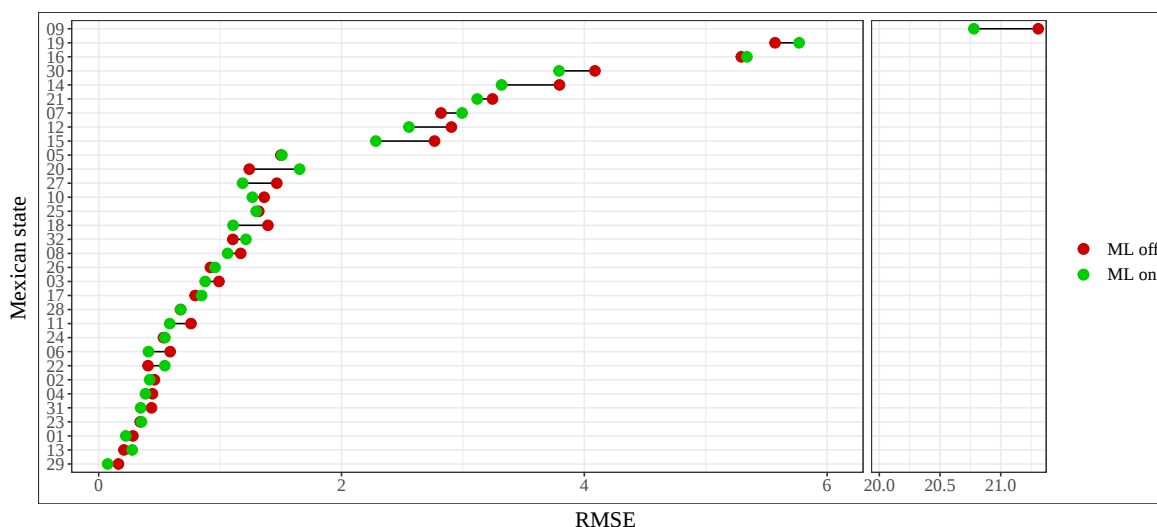


Figura. A.6: Error cuadrático medio (RMSE) de la recaudación de impuestos sobre la nómina prevista y real en 10 años simulados por estado con aprendizaje automático activado y desactivado. La ruptura del eje x se realizó para utilizar de manera efectiva el espacio de trazado y tratar los valores atípicos (Xu et al., 2021).

A.2.2. El efecto del aprendizaje automático en la simulación

Figura A.7 muestra las variaciones en los resultados de las tres variables observadas cuando se “enciende” o “apaga” el Aprendizaje Automático y se introducen cambios en la corrupción o inseguridad percibida. El valor observado en los datos corresponde a la coordenada 0,0. Los valores negativos representan una disminución en el valor percibido en el respectivo factor institucional, es decir, menos corrupción o inseguridad. Como se discutió anteriormente, sobre la base de la evidencia encontrada en los países en desarrollo, se espera que los determinantes institucionales como la confianza en el gobierno y la calidad de los servicios públicos influyan en la evasión fiscal (Daude et al., 2012).

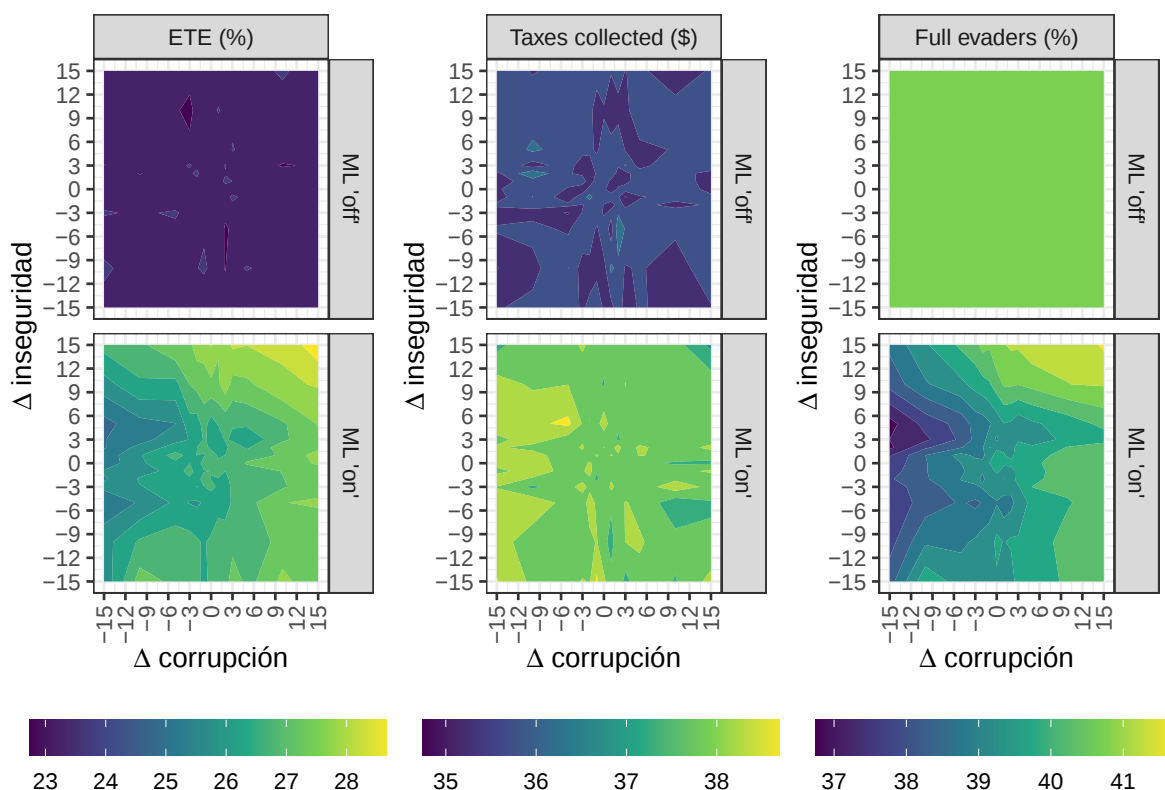


Figura. A.7: Porcentaje de extensión de la evasión fiscal, unidades monetarias de impuestos recaudados y porcentaje de evasores totales en el sistema, al variar la percepción sobre corrupción e inseguridad, con y sin aprendizaje automático. Las regiones violáceas denotan un valor pequeño de la variable de salida. Un valor bajo de ETE o evasores totales es bueno, mientras que para los impuestos recaudados, un valor bajo es malo.

Sin embargo, se aprecia que los resultados de la simulación con aprendizaje automático

“apagado” no capturan los cambios en los determinantes institucionales, principalmente en dos de las variables observadas, Extensión de la Evasión Tributaria y el número de evasores totales en el sistema. Esto significa que el modelo analítico da mayor importancia a los parámetros fiscales (tasa impositiva y tasa de penalización), y menos a los parámetros institucionales (variaciones en la percepción de corrupción e inseguridad).

Por el contrario, el modelo con aprendizaje “activado” puede reflejar las variaciones de los factores institucionales de forma plausible. Por ejemplo, si aumentan tanto la percepción de corrupción como la inseguridad, también aumenta la Extensión de la Evasión Fiscal. Paradójicamente, en ese punto también hay una alta recaudación, generada por el mayor número de sanciones cobradas. También se puede señalar que la percepción de corrupción juega un papel importante en la determinación de los empleadores para alinearse en el sector formal o informal. Pocos evasores en el sistema y alta recaudación se logran si la percepción de corrupción disminuye, incluso si la percepción de inseguridad aumenta levemente.

A.3. Conclusiones

En esta sección se describió de la mejor manera posible que permite el protocolo ODD, un modelo que combina herramientas aprendizaje automático y MBA. Tomando como caso de estudio el análisis de los factores institucionales sobre la evasión fiscal del Impuesto sobre la Nómina, fue posible formular un conjunto de observaciones que pueden ayudar a definir un mejor marco para la aplicación combinada de estas dos poderosas herramientas.

El Impuesto sobre la Nómina es un impuesto que en México tiene una tasa baja pero una multa alta por el monto evadido. Con estas características, los modelos analíticos no logran captar las diferencias producidas por los determinantes institucionales, que parecen tener menos importancia en la formulación matemática, pero se ha demostrado que, de hecho, sí tienen efecto. La inclusión de aprendizaje en los modelos basados en agentes permite a los agentes tomar decisiones que se aproximan mejor a la dinámica completa de los sistemas. Por lo tanto, una recomendación de política pública para influir en la decisión de los empleadores y ayudar a reducir el alcance de la evasión fiscal está relacionada con el fortalecimiento de las instituciones y la mejora de la calidad de los bienes públicos provistos.

Sin embargo, se señaló que los componentes actuales del protocolo ODD, incluido el aprendizaje en conceptos de diseño, no pueden describir formalmente la integración de



aprendizaje automático y MBA. Por lo tanto, el trabajo de investigación futuro consiste en la integración formal del aprendizaje automático en el marco del protocolo ODD, que debería servir como un conjunto de mejores prácticas o como asesoramiento y eventualmente formar la base para un conjunto más concreto de pautas para futuros desarrolladores de modelos. Esas pautas permitirán al desarrollador responder, ¿qué aprender? ¿Por qué aprender? Cuando aprender ¿Quién aprende y qué tipo de modelo de aprendizaje automático usar?

- Abdi, H. y Williams, L. J. (2010). Principal component analysis. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2(4):433–459.
- Abdulkareem, S., Augustijn, E.-W., Mustafa, Y., y Filatova, T. (2018). Intelligent judgements over health risks in a spatial agent-based model. *International Journal of Health Geographics*, 17(1).
- Abdulkareem, S., Mustafa, Y., Augustijn, E.-W., y Filatova, T. (2019). Bayesian networks for spatial learning: a workflow on using limited survey data for intelligent learning in spatial agent-based models. *Geoinformatica*, 23(2):243–268.
- Achter, S., Borit, M., Chattoe-Brown, E., y Siebers, P.-O. (2022). RAT-RS: a reporting standard for improving the documentation of data use in agent-based modelling. *International Journal of Social Research Methodology*, 25(4):517–540.
- Aghaie, A. y Heidary, M. (2019). Simulation-based optimization of a stochastic supply chain considering supplier disruption: Agent-based modeling and reinforcement learning. *Scientia Iranica*, 26(6):3780–3795.
- Aguirre-Quezada, J. P. y Sánchez-Ramírez, M. C. (2019). Evasión fiscal en México. Technical report, Instituto Belisario Domínguez.
- Akoglu, H. (2018). User's guide to correlation coefficients. *Turkish Journal of Emergency Medicine*, 18(3):91–93.
- Allingham, M. G. y Sandmo, A. (1972). Income tax evasion: a theoretical analysis. *Journal of Public Economics*, 1(3-4):323–338.

- Alm, J. (2011). Measuring, explaining, and controlling tax evasion: lessons from theory, experiments, and field studies. *International Tax and Public Finance*, 19(1):54–77.
- Alm, J. y Martinez-Vazquez, J. (2007). Tax morale and tax evasion in latin america. Technical report, International Studies Program.
- Anzola, D. (2019). Knowledge transfer in agent-based computational social science. *Studies in History and Philosophy of Science Part A*, 77:29–38.
- Aria, M. y Cuccurullo, C. (2017). bibliometrix: An r-tool for comprehensive science mapping analysis. *Journal of Informetrics*, 11(4):959–975.
- Ashreeta, P., Holzhauer, S., y Krebs, F. (2019). Overview of machine learning and data-driven methods in agent-based modeling of energy markets. In David K., Geihs K., L. M. y G., S., editors, *INFORMATIK 2019: 50 Jahre Gesellschaft für Informatik – Informatik für Gesellschaft*, volume 294, pages 571–584. Gesellschaft fur Informatik (GI).
- Augustijn, E.-W., Abdulkareem, S., Sadiq, M., y Albabawat, A. (2020). Machine learning to derive complex behaviour in agent-based modellzing. In *2020 International Conference on Computer Science and Software Engineering (CSASE)*, pages 284–289. Institute of Electrical and Electronics Engineers Inc.
- Axtell, R. L. (2000). Why agents? on the varied motivations for agent computing in the social sciences. In *Workshop on Agent Simulation: Applications, Models, and Tools*.
- Ayala-Gaytán, E. A., Cabral-Torres, R., y Chapa-Cantú, J. C. (2016). ¿quién paga el impuesto sobre nóminas de los estados de méxico?: Un análisis de incidencia económica. *Centro de Estudios de las Finanzas Públicas*, pages 1–43.
- Balasubramaniam, C. y Thigarasu, V. (2015). Agent-based modeling in supply chain management using improved c4-5. *Research Journal of Applied Sciences, Engineering and Technology*, 9(2):91–97.
- Baldwin, W. C., Sauser, B., y Cloutier, R. (2015). Simulation approaches for system of systems: Events-based versus agent based modeling. *Procedia Computer Science*, 44:363–372. 2015 Conference on Systems Engineering Research.
- Basilgan, M. y Christiansen, B. (2016). *Taxpayers' Attitudes towards Tax Evasion in Latin American Countries*. IGI Global.
- Bastian, M., Heymann, S., y Jacomy, M. (2009). Gephi: An open source software for exploring and manipulating networks.

- Batagelj, V. y Mrvar, A. (2002). Pajek— analysis and visualization of large networks. In *Graph Drawing*, pages 477–478. Springer Berlin Heidelberg.
- Batata, O., Augusto, V., y Xie, X. (2019). Mixed machine learning and agent-based simulation for respite care evaluation. In *2018 Winter Simulation Conference (WSC)*, volume 2018-December, pages 2668–2679. Institute of Electrical and Electronics Engineers Inc.
- Becker, G. S. (1968). Crime and punishment: An economic approach. *Journal of Political Economy*, 76(2):169–217.
- Bhattacharjee, S., MacPherson, B., Wang, R. F., y Gras, R. (2019). Animal communication of fear and safety related to foraging behavior and fitness: An individual-based modeling approach. *Ecological Informatics*, 54:101011.
- Billari, F. C., Ongaro, F., y Prskawetz, A. (2003). Introduction: Agent-based computational demography. In *Agent-Based Computational Demography*, pages 1–17. Physica-Verlag HD.
- Bloembergen, D., Tuyls, K., Hennes, D., y Kaisers, M. (2015). Evolutionary dynamics of multi-agent learning: A survey. *Journal of Artificial Intelligence Research*, 53:659–697.
- Bonet, J. y Rueda, F. (2012). Esfuerzo fiscal en los estados mexicanos. Technical report, Inter-American Development Bank.
- Borgonovo, E., Pangallo, M., Rivkin, J., Rizzo, L., y Siggelkow, N. (2022). Sensitivity analysis of agent-based models: a new protocol. *Computational and Mathematical Organization Theory*, 28(1):52–94.
- Breckling, B. (2002). Individual-based modelling potentials and limitations. *The Scientific World Journal*, 2:1044–1062.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1):5–32.
- Broniec, W., An, S., Rugarber, S., y Goel, A. K. (2021). Guiding parameter estimation of agent-based modeling through knowledge-based function approximation. In *Spring Symposium on Combining Machine Learning and Knowledge Engineering*.
- Broto, B., Bachoc, F., y Depecker, M. (2020). Variance reduction for estimation of shapley effects and adaptation to unknown input distribution. *SIAM/ASA Journal on Uncertainty Quantification*, 8(2):693–716.
- Caiani, A. y Caverzasi, E. (2017). Decentralized interacting macroeconomics and the agent-based “modellaccio”. In Gallegati, M., Palestini, A., y Russo, A., editors, *Introduction to*

- Agent-Based Economics*, pages 15–64. Academic Press.
- Candadai, M., Setzler, M., Izquierdo, E., y Froese, T. (2019). Embodied dyadic interaction increases complexity of neural dynamics: A minimal agent-based simulation model. *Frontiers in Psychology*, 10(MAR).
- Chakraborti, A. y Patriarca, M. (2008). Gamma-distribution and wealth inequality. *Pramana - Journal of Physics*, 71(2):233–243.
- Charteris, P., Golden, B., y Garrick, D. J. (2001). Livestock breeding industries as complex adaptive systems. In *Conference of the Association for the Advancement of Animal Breeding and Genetics*, volume 14, pages 461–464.
- Chen, S., Londoño-Larrea, P., McGough, A., Bible, A., Gunaratne, C., Araujo-Granda, P., Morrell-Falvey, J., Bhowmik, D., y Fuentes-Cabrera, M. (2021). Application of machine learning techniques to an agent-based model of pantoea. *Frontiers in Microbiology*, 12.
- Chow, C. y Liu, C. (1968). Approximating discrete probability distributions with dependence trees. *IEEE Trans. Inf. Theory*, 14:462–467.
- Chu, Z., Yang, B., Ha, C., y Ahn, K. (2018). Modeling gdp fluctuations with agent-based model. *Physica A: Statistical Mechanics and its Applications*, 503:572–581.
- Clouvel, L., looss, B., Chabridon, V., Il Idrissi, M., y Robin, F. (2023). Variance-based importance measures in the linear regression context: Review, new insights and applications. working paper or preprint.
- Cockrell, C. y An, G. (2021). Utilizing the heterogeneity of clinical data for model refinement and rule discovery through the application of genetic algorithms to calibrate a high-dimensional agent-based model of systemic inflammation. *Frontiers in Physiology*, 12.
- CONAPO (2020). Datos abiertos. indicadores demográficos 1950 - 2050.
- Conte, R. y Paolucci, M. (2014). On agent-based modeling and computational social science. *Frontiers in Psychology*, 5:668.
- Conte, R. y Terna, P. (2000). Una discussione sulla simulazione in campo sociale: mente e società. *Sistemi intelligenti*, page 321–346.
- Cummings, P. y Crooks, A. (2020). Development of a hybrid machine learning agent based model for optimization and interpretability. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 12268 LNCS:151–160.

- Dalle, O. (2012). On reproducibility and traceability of simulations. In *Proceedings Title: Proceedings of the 2012 Winter Simulation Conference (WSC)*. IEEE.
- Daude, C., Gutiérrez, H., y Ángel Melguizo (2012). What drives tax morale?
- Davidsson, P. (2001). Categories of artificial societies. In *Engineering Societies in the Agents World II*, pages 1–9. Springer Berlin Heidelberg.
- Davis, J. S., Hecht, G., y Perkins, J. D. (2003). Social behaviors, enforcement, and tax compliance dynamics. *The Accounting Review*, 78(1):39–69.
- Day, C. C., Zollner, P. A., Gilbert, J. H., y McCann, N. P. (2020). Individual-based modeling highlights the importance of mortality and landscape structure in measures of functional connectivity. *Landscape Ecology*, 35(10):2191–2208.
- Dehghanpour, K., Hashem Nehrir, M., Sheppard, J., y Kelly, N. (2018). Agent-based modeling of retail electrical energy markets with demand response. *IEEE Transactions on Smart Grid*, 9(4):3465–3475.
- Dehghanpour, K., Nehrir, M., Sheppard, J., y Kelly, N. (2016). Agent-based modeling in electrical energy markets using dynamic bayesian networks. *IEEE Transactions on Power Systems*, 31(6):4744–4754.
- Dessalles, J. L., Ferber, J., y Phan, D. (2008). Emergence in agent-based computational social science. In *Intelligent Complex Adaptive Systems*, pages 255–299. IGI Global.
- Dignum, V., Gilbert, N., y Wellman, M. P. (2016). Introduction to the special issue on autonomous agents for agent-based modeling. *Autonomous Agents and Multi-Agent Systems*, 30(6):1021–1022.
- Doran, J. (1998). Social simulation, agents and artificial societies. In *Proceedings International Conference on Multi Agent Systems (Cat. No.98EX160)*, pages 4–5.
- Edali, M. y Yücel, G. (2019). Exploring the behavior space of agent-based simulation models using random forest metamodels and sequential sampling. *Simulation Modelling Practice and Theory*, 92:62–81.
- Edmonds, B. (2000). Complexity and scientific modelling. *Foundations of Science*, 5(3):379–390.
- Edmonds, B., Page, C. L., Bithell, M., Chattoe-Brown, E., Grimm, V., Meyer, R., Montañola-Sales, C., Ormerod, P., Root, H., y Squazzoni, F. (2019). Different modelling purposes. *Journal of Artificial Societies and Social Simulation*, 22(3).

- Efron, S. y Ravid, R. (2018). *Writing the Literature Review: A Practical Guide*. Guilford Publications.
- Epstein, J. M. (2008). Why model? *Journal of Artificial Societies and Social Simulation*, 11(4):12.
- Epstein, J. M. y Axtell, R. L. (1997). *Growing Artificial Societies: Social Science from the Bottom Up*. MIT Press, USA, 1st edition.
- Esmaeili Aliabadi, D., Kaya, M., y Sahin, G. (2017). Competition, risk and learning in electricity markets: An agent-based simulation study. *Applied Energy*, 195:1000–1011.
- Fagiolo, G., Guerini, M., Lamperti, F., Moneta, A., y Roventini, A. (2019). Validation of agent-based models in economics and finance. In *Simulation Foundations, Methods and Applications*, pages 763–787. Springer International Publishing.
- Fievet, L. y Sornette, D. (2018). Calibrating emergent phenomena in stock markets with agent based models. *PLoS ONE*, 13(3).
- Fischer, G. (1991). The importance of models in making complex systems comprehensible. In *Mental Models and Human-Computer Interaction 2*, pages 3–36. Elsevier.
- Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *The Annals of Statistics*, 29(5):1189–1232.
- Friedman, N., Geiger, D., y Goldszmidt, M. (1997). Bayesian network classifiers. *Machine Learning*, 29(2/3):131–163.
- Fuller, D., de Arruda, E., y Ferreira Filho, V. (2020). Learning-agent-based simulation for queue network systems. *Journal of the Operational Research Society*, 71(11):1723–1739.
- García-Magariño, I., Medrano, C., y Delgado, J. (2020). Estimation of missing prices in real-estate market agent-based simulations with machine learning and dimensionality reduction methods. *Neural Computing and Applications*, 32(7):2665–2682.
- Garfield, E., Paris, S. W., y Stock, W. G. (2006). Histcite™ : A software tool for informetric analysis of citation linkage. *Information-Wissenschaft und Praxis*, 57:391–400.
- Garg, A., Yuen, S., Seekhao, N., Yu, G., Karwowski, J., Powell, M., Sakata, J., Mongeau, L., JaJa, J., y Li-Jessen, N. (2019). Towards a physiological scale of vocal fold agent-based models of surgical injury and repair: Sensitivity analysis, calibration and verification. *Applied Sciences (Switzerland)*, 9(15).
- Gass, S. I. (1984). Documenting a computer-based model. *Interfaces*, 14(3):84–93.

- Gavin, M. (2018). An agent-based computational approach to "the adam smith problem". *Historical Social Research*, 43(1):308–336.
- Ghaffarian, S., Roy, D., Filatova, T., y Kerle, N. (2021). Agent-based modelling of post-disaster recovery with remote sensing data. *International Journal of Disaster Risk Reduction*, 60.
- Ghorbani, A., Dechesne, F., Dignum, V., y Jonker, C. (2014). Enhancing ABM into an inevitable tool for policy analysis. *Journal on Policy and Complex Systems*, 1(1):61–76.
- Giere, R. N. (1999). Using models to represent reality. In *Model-Based Reasoning in Scientific Discovery*, pages 41–57. Springer US.
- Gilbert, G. N. (1997). Environments and languages to support social simulation: Summary of an informal discussion. *Bulletin of Sociological Methodology*, 56(1):79–80.
- Gilbert, N. (2020). *Agent-Based Models*. SAGE Publications, Inc.
- Goodell, J. W., Kumar, S., Lim, W. M., y Pattnaik, D. (2021). Artificial intelligence and machine learning in finance: Identifying foundations, themes, and research clusters from bibliometric analysis. *Journal of Behavioral and Experimental Finance*, 32:100577.
- Gościk, J. y Łukaszuk, T. (2013). Application of the recursive feature elimination and the relaxed linear separability feature selection algorithms to gene expression data analysis. *Advances in Computer Science Research*, page 39–52.
- Grimm, V. (1999). Ten years of individual-based modelling in ecology: what have we learned and what could we learn in the future? *Ecological Modelling*, 115(2-3):129–148.
- Grimm, V., Berger, U., Bastiansen, F., Eliassen, S., Ginot, V., Giske, J., Goss-Custard, J., Grand, T., Heinz, S. K., Huse, G., Huth, A., Jepsen, J. U., Jørgensen, C., Mooij, W. M., Müller, B., Pe'er, G., Piou, C., Railsback, S. F., Robbins, A. M., Robbins, M. M., Rossmanith, E., Rüger, N., Strand, E., Souissi, S., Stillman, R. A., Vabø, R., Visser, U., y DeAngelis, D. L. (2006). A standard protocol for describing individual-based and agent-based models. *Ecological Modelling*, 198(1-2):115–126.
- Grimm, V., Berger, U., DeAngelis, D. L., Polhill, J. G., Giske, J., y Railsback, S. F. (2010). The odd protocol: A review and first update. *Ecological Modelling*, 221(23):2760–2768.
- Grimm, V., Railsback, S. F., Vincenot, C. E., Berger, U., Gallagher, C., DeAngelis, D. L., Edmonds, B., Ge, J., Giske, J., Groeneveld, J., Johnston, A. S., Milles, A., Nabe-Nielsen, J., Polhill, G., Radchuk, V., Rohwäder, M.-S., Stillman, R. A., Thiele, J. C., y Ayllón, D. (2020). The odd protocol for describing agent-based and other simulation models: A second update

- to improve clarity, replication, and structural realism. *Journal of Artificial Societies and Social Simulation*, 23(2):7.
- Guerini, M. y Moneta, A. (2017). A method for agent-based models validation. *Journal of Economic Dynamics and Control*, 82:125–141.
- Gursoy, F. y Badur, B. (2021). An agent-based modeling approach to brain drain. *IEEE Transactions on Computational Social Systems*.
- Han, S. y Kim, H. (2021). Optimal feature set size in random forest regression. *Applied Sciences*, 11(8).
- Harati, S., Perez, L., y Molowny-Horas, R. (2021). Promoting the emergence of behavior norms in a principal–agent problem—an agent-based modeling approach using reinforcement learning. *Applied Sciences (Switzerland)*, 11(18).
- Hayashi, S., Prasasti, N., Kanamori, K., y Ohwada, H. (2016). Improving behavior prediction accuracy by using machine learning for agent-based simulation. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 9621:280–289.
- Haykin, S., Haykin, S., y HAYKIN, S. (1999). *Neural Networks: A Comprehensive Foundation*. International edition. Prentice Hall.
- Hearst, M., Dumais, S., Osuna, E., Platt, J., y Scholkopf, B. (1998). Support vector machines. *IEEE Intelligent Systems and their Applications*, 13(4):18–28.
- Heath, B. L. y Hill, R. R. (2008). The early history of agent-based modeling. *IIE Annual Conference.Proceedings*, pages 971–976. Copyright - Copyright Institute of Industrial Engineers-Publisher 2008; Última actualización - 2021-09-09.
- Hoi, S. C., Sahoo, D., Lu, J., y Zhao, P. (2021). Online learning: A comprehensive survey. *Neurocomputing*, 459:249–289.
- Hokamp, S. (2014). Dynamics of tax evasion with back auditing, social norm updating, and public goods provision – An agent-based simulation. *Journal of Economic Psychology*, 40(C):187–199.
- Husereau, D., Drummond, M., Petrou, S., Carswell, C., Moher, D., Greenberg, D., Augustovski, F., Briggs, A. H., Mauskopf, J., y and, E. L. (2013). Consolidated health economic evaluation reporting standards (CHEERS) statement. *BMJ*, 346(mar25 1):f1049–f1049.
- Hyun, J.-Y., Huang, S.-Y., Yang, Y.-C., Tidwell, V., y Macknick, J. (2019). Using a coupled

- agent-based modeling approach to analyze the role of risk perception in water management decisions. *Hydrology and Earth System Sciences*, 23(5):2261–2278.
- Ibrahim, M., Hashmi, U., Nabeel, M., Imran, A., y Ekin, S. (2021). Embracing complexity: Agent-based modeling for hetnets design and optimization via concurrent reinforcement learning algorithms. *IEEE Transactions on Network and Service Management*.
- INEGI (2019a). Encuesta nacional de ocupación y empleo (enoe). <https://www.inegi.org.mx/programas/enoe/15ymas/>.
- INEGI (2019b). Finanzas públicas estatales y municipales.
- INEGI (2019c). Medición de la economía informal. <https://www.inegi.org.mx/temas/pibmed/>.
- Jalalimanesh, A., Haghighi, H., Ahmadi, A., Hejazian, H., y Soltani, M. (2017a). Multi-objective optimization of radiotherapy: distributed q-learning and agent-based simulation. *Journal of Experimental and Theoretical Artificial Intelligence*, 29(5):1071–1086.
- Jalalimanesh, A., Shahabi Haghighi, H., Ahmadi, A., y Soltani, M. (2017b). Simulation-based optimization of radiotherapy: Agent-based modeling and reinforcement learning. *Mathematics and Computers in Simulation*, 133:235–248.
- Janssen, S., Sharpanskykh, A., Curran, R., y Langendoen, K. (2019). Using causal discovery to analyze emergence in agent-based models. *Simulation Modelling Practice and Theory*, 96.
- Joanes, D. N. y Gill, C. A. (1998). Comparing measures of sample skewness and kurtosis. *Journal of the Royal Statistical Society. Series D (The Statistician)*, 47(1):183–189.
- Johnson, D. H. (2001). Validating and evaluating models. In *Modeling in Natural Resource Management: Development, Interpretation, and Application*. Island Press, Washington, DC.
- Jørgensen, A. C. S., Ghosh, A., Sturrock, M., y Shahrezaei, V. (2022). Efficient bayesian inference for stochastic agent-based models. *PLOS Computational Biology*, 18(10):e1009508.
- Judson, O. P. (1994). The rise of the individual-based model in ecology. *Trends in Ecology and Evolution*, 9(1):9–14.
- Jäger, G. (2019). Replacing rules by neural networks a framework for agent-based modelling. *Big Data and Cognitive Computing*, 3(4):1–12.
- Jäger, G. (2021). Using neural networks for a universal framework for agent-based models. *Mathematical and Computer Modelling of Dynamical Systems*, 27(1):162–178.

- Kavak, H., Padilla, J., Lynch, C., y Diallo, S. (2018). Big data, agents, and machine learning: Towards a data-driven agent-based modeling approach. In Frydenlund E., Shafer S., K. H., editor, *ANSS 201818: Proceedings of the Annual Simulation Symposium*, volume 50, pages 125–136. The Society for Modeling and Simulation International.
- Khlif, H. y Achek, I. (2015). The determinants of tax evasion: a literature review. *International Journal of Law and Management*, pages 486–497.
- Kim, D., Yun, T.-S., Moon, I.-C., y Bae, J. (2021). Automatic calibration of dynamic and heterogeneous parameters in agent-based models. *Autonomous Agents and Multi-Agent Systems*, 35(2).
- Klein, D., Marx, J., y Fischbach, K. (2018). Agent-based modeling in social science, history, and philosophy. an introduction. *Historical Social Research / Historische Sozialforschung Vol. 43*, No. 1: Volumes per year: 1
- Koda, H., Arai, Z., y Matsuda, I. (2020). Agent-based simulation for reconstructing social structure by observing collective movements with special reference to single-file movement. *PLoS ONE*, 15(12 December).
- Kostadinov, F., Holm, S., Steubing, B., Thees, O., y Lemm, R. (2014). Simulation of a swiss wood fuel and roundwood market: An explorative study in agent-based modeling. *Forest Policy and Economics*, 38:105–118.
- Krivorotko, O., Sosnovskaia, M., Vashchenko, I., Kerr, C., y Lesnic, D. (2022). Agent-based modeling of COVID-19 outbreaks for new york state and UK: Parameter identification algorithm. *Infectious Disease Modelling*, 7(1):30–44.
- Kuhn, M. (2008). Building predictive models in r using the caret package. *Journal of Statistical Software*, 28(5).
- Lakić, E., Artač, G., y Gubina, A. (2015). Agent-based modeling of the demand-side system reserve provision. *Electric Power Systems Research*, 124:85–91.
- Lamperti, F., Roventini, A., y Sani, A. (2018). Agent-based model calibration using machine learning surrogates. *Journal of Economic Dynamics and Control*, 90:366–389.
- Larie, D., An, G., y Cockrell, R. (2021). The use of artificial neural networks to forecast the behavior of agent-based models of pathophysiology: An example utilizing an agent-based model of sepsis. *Frontiers in Physiology*, 12.
- Lee, J.-S., Filatova, T., Ligmann-Zielinska, A., Hassani-Mahmooui, B., Stonedahl, F., Lorscheid,

- I., Voinov, A., Polhill, J. G., Sun, Z., y Parker, D. C. (2015). The complexities of agent-based modeling output analysis. *Journal of Artificial Societies and Social Simulation*, 18(4):4.
- Lee, M. y Hong, T. (2019). Hybrid agent-based modeling of rooftop solar photovoltaic adoption by integrating the geographic information system and data mining technique. *Energy Conversion and Management*, 183:266–279.
- Lee, T.-C. y Wong, K. (2016). An agent-based model for queue formation of powered two-wheelers in heterogeneous traffic. *Physica A: Statistical Mechanics and its Applications*, 461:199–216.
- Leibo, J. Z., Zambaldi, V., Lanctot, M., Marecki, J., y Graepel, T. (2017). Multi-agent reinforcement learning in sequential social dilemmas. In *Proceedings of the 16th Conference on Autonomous Agents and MultiAgent Systems*, AAMAS '17, page 464–473, Richland, SC. International Foundation for Autonomous Agents and Multiagent Systems.
- Lengnick, M. (2013). Agent-based macroeconomics: A baseline model. *Journal of Economic Behavior & Organization*, 86:102–120.
- Li, X., Mao, W., Zeng, D., y Wang, F.-Y. (2008). Agent-based social simulation and modeling in social computing. In *Intelligence and Security Informatics*, pages 401–412. Springer Berlin Heidelberg.
- Liang, Y., Guo, C., Ding, Z., y Hua, H. (2020). Agent-based modeling in electricity market using deep deterministic policy gradient algorithm. *IEEE Transactions on Power Systems*, 35(6):4180–4192.
- Ligmann, A. y Sun, L. (2010). Applying time-dependent variance-based global sensitivity analysis to represent the dynamics of an agent-based model of land use change. *International Journal of Geographical Information Science*, 24:1829 – 1850.
- Lisic, J. J. y Cruze, N. B. (2016). Local pivotal methods for large surveys. In *International Conference on Establishment Surveys*.
- Ma, P., Han, X.-H., Lin, Y., Moore, J., Guo, Y.-X., y Yue, M. (2019). Exploring the relative importance of biotic and abiotic factors that alter the self-thinning rule: Insights from individual-based modelling and machine-learning. *Ecological Modelling*, 397:16–24.
- MacPherson, B., Mashayekhi, M., Gras, R., y Scott, R. (2017). Exploring the connection between emergent animal personality and fitness using a novel individual-based model and decision tree approach. *Ecological Informatics*, 40:81–92.

- Malleson, N., Heppenstall, A., y See, L. (2010). Crime reduction through simulation: An agent-based model of burglary. *Computers, Environment and Urban Systems*, 34(3):236–250.
- Manzo, G. (2014). Potentialités et limites de la simulation multi-agents : une introduction. *Revue française de sociologie*, 55(4):653.
- Marks, R. E. (2013). Validation and model selection: Three similarity measures compared. *Complexity Economics*, 2(1):41–61.
- Marzouk, M. y Hassan, F. (2022). Modeling evacuation and visitation proximity in museums using agent-based simulation. *Journal of Building Engineering*, 56:104794.
- Masuda, S., Bahr, K., Tsuchiya, N., y Takemori, T. (2022). Agent based simulation with data driven parameterization for evaluation of social acceptance of a geothermal development: a case study in tsuchiya, fukushima, japan. *Scientific Reports*, 12(1).
- Michalski, R. S., Carbonell, J. G., y Mitchell, T. M., editors (1983). *Machine Learning*. Springer Berlin Heidelberg.
- Miller, M. Z., Griendling, K., y Mavris, D. N. (2012). Exploring human factors effects in the smart grid system of systems demand response. In *2012 7th International Conference on System of Systems Engineering (SoSE)*, pages 1–6.
- Mittone, L. y Patelli, P. (2000). Imitative behaviour in tax evasion. In *Economic Simulations in Swarm: Agent-Based Modelling and Object Oriented Programming*, pages 133–158. Springer US.
- Mnih, V., Kavukcuoglu, K., Silver, D., Graves, A., Antonoglou, I., Wierstra, D., y Riedmiller, M. A. (2013). Playing atari with deep reinforcement learning. *CoRR*, abs/1312.5602.
- Monks, T., Currie, C. S. M., Onggo, B. S., Robinson, S., Kunc, M., y Taylor, S. J. E. (2018). Strengthening the reporting of empirical simulation studies: Introducing the STRESS guidelines. *Journal of Simulation*, 13(1):55–67.
- Montgomery, D., Peck, E., y Vining, G. (2015). *Introduction to Linear Regression Analysis*. Wiley Series in Probability and Statistics. Wiley.
- Navarro, J. F., Frenk, C. S., y White, S. D. M. (1997). A universal density profile from hierarchical clustering. *The Astrophysical Journal*, 490(2):493–508.
- Neri, F. (2019). Combining machine learning and agent based modeling for gold price prediction. *Communications in Computer and Information Science*, 900:91–100.
- Norman, M., Koehler, M., Kutarnia, J., Silvey, P., Tolk, A., y Tracy, B. (2018). Applying com-

- plexity science with machine learning, agent-based models, and game engines: Towards embodied complex systems engineering. In *Unifying Themes in Complex Systems IX*, pages 173–183. Springer.
- Nourqolipour, R. y Shariff, R. (2010). How agent based modeling (abm) can be linked to gis for modelling land use and land cover change. In *MRSS 6th International Remote Sensing and GIS Conference and Exhibition*.
- Nunes, A., Zwick, M., y Wakeland, W. (2021). Sensitivity analysis of an agent-based simulation model using reconstructability analysis. *International Journal of General Systems*, 50(3):319–338.
- Omidshafiei, S., Pazis, J., Amato, C., How, J. P., y Vian, J. (2017). Deep decentralized multi-task multi-agent reinforcement learning under partial observability. In *Proceedings of the 34th International Conference on Machine Learning - Volume 70*, ICML'17, page 2681–2690. JMLR.
- Ozik, J., Collier, N., Wozniak, J., Macal, C., y An, G. (2018). Extreme-scale dynamic exploration of a distributed agent-based model with the emews framework. *IEEE Transactions on Computational Social Systems*, 5(3):884–895.
- Papadopoulos, S. y Azar, E. (2016). Integrating building performance simulation in agent-based modeling using regression surrogate models: A novel human-in-the-loop energy modeling approach. *Energy and Buildings*, 128:214–223.
- Pawar, S. y Jha, A. K. (2022). Analysis of residential location choices of different socio-economic groups and their impact on the density in a city using agent based modelling. *Applied Spatial Analysis and Policy*, 16(1):119–139.
- Pearce, P. y Slade, R. (2018). Feed-in tariffs for solar microgeneration: Policy evaluation and capacity projections using a realistic agent-based model. *Energy Policy*, 116:95–111.
- Perry, G. y O'Sullivan, D. (2018). Identifying narrative descriptions in agent-based models representing past human-environment interactions. *Journal of Archaeological Method and Theory*, 25(3):795–817.
- Persson, O. (2008). Bibexcel: A toolbox for bibliometricians. *International Society for Scientometrics and Informetrics*.
- Perumal, R. y Zyl, T. (2022). Surrogate-assisted strategies: the parameterisation of an infectious disease agent-based model. *Neural Computing and Applications*.

- Peters, R., Lin, Y., y Berger, U. (2016). Machine learning meets individual-based modelling: Self-organising feature maps for the analysis of below-ground competition among plants. *Ecological Modelling*, 326:142–151.
- Pinder, J. E., Wiener, J. G., y Smith, M. H. (1978). The weibull distribution: A new method of summarizing survivorship data. *Ecology*, 59(1):175–179.
- Platas-López, A., Guerra-Hernández, A., Grimaldo, F., Cruz-Ramírez, N., Mezura-Montes, E., y Quiroz-Castellanos, M. (2023a). dplbnDE: An r package for discriminative parameter learning of bayesian networks by differential evolution. *SoftwareX*, 23:101442.
- Platas-López, A., Guerra-Hernández, A., Quiroz-Castellanos, M., y Cruz-Ramírez, N. (2023b). A survey on agent-based modelling assisted by machine learning. *Expert Systems*.
- Platas-López, A., Mezura-Montes, E., Cruz-Ramírez, N., y Guerra-Hernández, A. (2021). Discriminative learning of bayesian network parameters by differential evolution. *Applied Mathematical Modelling*, 93:244–256.
- Platas-López, A. (2014). Esfuerzo fiscal en la recaudación de impuesto sobre nómina en México. *Horizontes de la Contaduría*, pages 89–103.
- Platas-López, A., Guerra-Hernández, A., Quiroz-Castellanos, M., y Cruz-Ramírez, N. (2023). Agent-based models assisted by supervised learning: A proposal for model specification. *Electronics*, 12(3).
- Pouladi, P., Afshar, A., Molajou, A., y Afshar, M. (2020). Socio-hydrological framework for investigating farmers' activities affecting the shrinkage of urmia lake; hybrid data mining and agent-based modelling. *Hydrological Sciences Journal*, 65(8):1249–1261.
- Price, K. y Storn, R. (1996). Minimizing the real functions of the icec'96 contest by differential evolution. In *in Proc. of IEEE C. Evol. Computat.*, pages 842–844.
- Pritchard, A. et al. (1969). Statistical bibliography or bibliometrics. *Journal of documentation*, 25(4):348–349.
- Quinlan, J. R. (1986). Induction of decision trees. *Machine Learning*, 1(1):81–106.
- R Core Team (2021). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria.
- R. Vahdati, A., Weissmann, J., Timmermann, A., Ponce de León, M., y Zollikofer, C. (2019). Drivers of late pleistocene human survival and dispersal: an agent-based modeling and machine learning approach. *Quaternary Science Reviews*, 221.

- Rahmandad, H. y Sterman, J. D. (2012). Reporting guidelines for simulation-based research in social sciences. *System Dynamics Review*, 28(4):396–411.
- Railsback, S. F. y Grimm, V. (2019). *Agent-Based and Individual-Based Modeling*. Princeton University Press, 41 William Street, Princeton, New Jersey, USA.
- Rajabi, M., Pilesjö, P., Shirzadi, M., Fadaei, R., y Mansourian, A. (2016). A spatially explicit agent-based modeling approach for the spread of cutaneous leishmaniasis disease in central iran, isfahan. *Environmental Modelling and Software*, 82:330–346.
- Rand, W. y Rust, R. T. (2011). Agent-based modeling in marketing: Guidelines for rigor. *International Journal of Research in Marketing*, 28(3):181–193.
- Reiker, T., Golumbeanu, M., Shattock, A., Burgert, L., Smith, T. A., Filippi, S., Cameron, E., y Penny, M. A. (2021). Emulator-based bayesian optimization for efficient multi-objective calibration of an individual-based model of malaria. *Nature Communications*, 12(1).
- Retzlaff, C. O., Ziefle, M., y Valdez, A. C. (2021). The history of agent-based modeling in the social sciences. In *Digital Human Modeling and Applications in Health, Safety, Ergonomics and Risk Management*, pages 304–319. Springer International Publishing.
- Romero-Mujalli, D., Cappelletto, J., Herrera, E., y Tárano, Z. (2017). The effect of social learning in a small population facing environmental change: an agent-based simulation. *Journal of Ethology*, 35(1):61–73.
- Rosés, R., Kadar, C., y Malleson, N. (2021). A data-driven agent-based simulation to predict crime patterns in an urban environment. *Computers, Environment and Urban Systems*, 89.
- Rummery, G. A. y Niranjan, M. (1994). On-line Q-learning using connectionist systems. Technical Report 166, Cambridge University Engineering Department.
- Russell, S., Russell, S., y Norvig, P. (2020). *Artificial Intelligence: A Modern Approach*. Pearson series in artificial intelligence. Pearson.
- Sadeghipour, M., Khoshnevisan, M., Jafari, A., y Shariatpanahi, S. (2017). Friendship network and dental brushing behavior among middle school students: An agent based modeling approach. *PLoS ONE*, 12(1).
- Saeedi, S. (2018). Integrating macro and micro scale approaches in the agent-based modeling of residential dynamics. *International Journal of Applied Earth Observation and Geoinformation*, 68:214–229.
- Salle, I. (2015). Modeling expectations in agent-based models - an application to central

- bank's communication and monetary policy. *Economic Modelling*, 46:130–141.
- Saltelli, A., Annoni, P., Azzini, I., Campolongo, F., Ratto, M., y Tarantola, S. (2010). Variance based sensitivity analysis of model output. design and estimator for the total sensitivity index. *Computer Physics Communications*, 181(2):259–270.
- Saltelli, A., Ratto, M., Andres, T., Campolongo, F., Cariboni, J., Gatelli, D., Saisana, M., y Tarantola, S. (2007). *Global Sensitivity Analysis. The Primer*. Wiley.
- Samaniego Breach, N. (2014). La participación del trabajo en el ingreso nacional: el regreso a un tema olvidado. *Economía UNAM*, 11:52 – 77.
- Samuel, A. L. (1959). Some studies in machine learning using the game of checkers. *IBM Journal of Research and Development*, pages 71–105.
- Sandmo, A. (2005). The theory of tax evasion: A retrospective view. *National Tax Journal*, 58(4):643–663.
- Scanagatta, M., Salmerón, A., y Stella, F. (2019). A survey on bayesian network structure learning from data. *Prog. Artif. Intell.*, 8:425–439.
- Schauder, S., Thomsen, M., y Nayga Jr, R. (2020). Agent-based modeling insights into the optimal distribution of the fresh fruit and vegetable program. *Preventive Medicine Reports*, 20.
- Schelling, T. C. (1969). Models of segregation. *The American Economic Review*, 59(2):488–493.
- Schelling, T. C. (1971). Dynamic models of segregation†. *The Journal of Mathematical Sociology*, 1(2):143–186.
- Schmolke, A., Thorbek, P., DeAngelis, D. L., y Grimm, V. (2010). Ecological models supporting environmental decision making: a strategy for the future. *Trends in Ecology and Evolution*, 25(8):479–486.
- Schwabish, J. (2021). *Better Data Visualizations: A Guide for Scholars, Researchers, and Wonks*. Columbia University Press.
- Schwarz, C. V., Reiser, B. J., Davis, E. A., Kenyon, L., Achér, A., Fortus, D., Shwartz, Y., Hug, B., y Krajcik, J. (2009). Developing a learning progression for scientific modeling: Making scientific modeling accessible and meaningful for learners. *Journal of Research in Science Teaching*, 46(6):632–654.
- Scott, R., MacPherson, B., y Gras, R. (2018). A comparison of stable and fluctuating resources

- with respect to evolutionary adaptation and life-history traits using individual-based modeling and machine learning. *Journal of Theoretical Biology*, 459:52–66.
- Sculley, D. (2010). Web-scale k-means clustering. In *Proceedings of the 19th international conference on World wide web - WWW*. ACM Press.
- Sert, E., Bar-Yam, Y., y Morales, A. (2020). Segregation dynamics with reinforcement learning and agent based modeling. *Scientific Reports*, 10(1).
- Shiono, T. (2021). Estimation of agent-based models using bayesian deep learning approach of bayesflow. *Journal of Economic Dynamics and Control*, 125.
- Siegfried, R. (2014). *Modeling and Simulation of Complex Systems*. Springer Fachmedien Wiesbaden.
- Silver, D., Huang, A., Maddison, C. J., Guez, A., Sifre, L., van den Driessche, G., Schrittwieser, J., Antonoglou, I., Panneershelvam, V., Lanctot, M., Dieleman, S., Grewe, D., Nham, J., Kalchbrenner, N., Sutskever, I., Lillicrap, T., Leach, M., Kavukcuoglu, K., Graepel, T., y Hassabis, D. (2016). Mastering the game of go with deep neural networks and tree search. *nature*, 529(7587):484–489.
- Skinner, J. y Slemrod, J. (1985). An economic perspective on tax evasion. *National Tax Journal*, 38(3):345–353.
- Slemrod, J. y Yitzhaki, S. (2000). Tax avoidance, evasion, and administration. Technical report, National Bureau of Economic Research.
- Squazzoni, F. (2012). *Agent-Based Computational Sociology*. John Wiley & Sons, Ltd.
- Srinivasan, S., Lanctot, M., Zambaldi, V., Pérolat, J., Tuyls, K., Munos, R., y Bowling, M. (2018). Actor-critic policy optimization in partially observable multiagent environments. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, NIPS'18, page 3426–3439, Red Hook, NY, USA. Curran Associates Inc.
- Sutton, R. S. y Barto, A. G. (2018). *Reinforcement Learning*. The MIT Press, Cambridge, Massachusetts London, England.
- Sánchez-Marño, N., Alonso-Betanzos, A., Fontenla-Romero, O., Brinquis-Núñez, C., Polhill, J., Craig, T., Dumitru, A., y García-Mira, R. (2015). An agent-based model for simulating environmental behavior in an educational organization. *Neural Processing Letters*, 42(1):89–118.
- Sánchez-Marño, N., Alonso-Betanzos, A., Fontenla-Romero, O., Polhill, J., y Craig, T. (2017).

- Empirically-derived behavioral rules in agent-based models using decision trees learned from questionnaire data. *Understanding Complex Systems*, pages 53–76.
- Tanabe, R. y Fukunaga, A. (2014). Improving the search performance of shade using linear population size reduction. In *in Proc. of IEEE C. Evol. Computat.*, pages 1658–1665.
- Tarvid, A. (2016). Complex adaptive systems and agent-based modelling. In *Agent-Based Modelling of Social Networks in Labour–Education Market System*, pages 23–38.
- ten Broeke, G., van Voorn, G., Ligtenberg, A., y Molenaar, J. (2021). The use of surrogate models to analyse agent-based models. *Journal of Artificial Societies and Social Simulation*, 24(2):3.
- Tesfatsion, L. (2002). Agent-based computational economics: Growing economies from the bottom up. *Artificial Life*, 8(1):55–82.
- Thiele, J. C., Kurth, W., y Grimm, V. (2014). Facilitating parameter estimation and sensitivity analysis of agent-based models: A cookbook using NetLogo and 'r'. *Journal of Artificial Societies and Social Simulation*, 17(3).
- Tieleman, S. (2021). Towards a validation methodology for macroeconomic agent-based models. *Computational Economics*, 60(4):1507–1527.
- Uhrmacher, A. M. y Weyns, D. (2009). *Multi-Agent Systems: Simulation and Applications*. CRC Press, Inc., Boca Raton, FL, USA.
- van der Hoog, S. (2019). Surrogate modelling in (and of) agent-based models: A prospectus. *Computational Economics*, 53(3):1245–1263.
- van Eck, N. J. y Waltman, L. (2010). Software survey: Vosviewer, a computer program for bibliometric mapping. *Scientometrics*, 84(2):523–538.
- Waltemath, D., Adams, R., Beard, D. A., Bergmann, F. T., Bhalla, U. S., Britten, R., Chelliah, V., Cooling, M. T., Cooper, J., Crampin, E. J., Garny, A., Hoops, S., Hucka, M., Hunter, P., Klipp, E., Laibe, C., Miller, A. K., Moraru, I., Nickerson, D., Nielsen, P., Nikolski, M., Sahle, S., Sauro, H. M., Schmidt, H., Snoep, J. L., Tolle, D., Wolkenhauer, O., y Novère, N. L. (2011). Minimum information about a simulation experiment (MIASE). *PLoS Computational Biology*, 7(4):e1001122.
- Wang, C., Hu, M., Yang, L., y Zhao, Z. (2021). Prediction of air traffic delays: An agent-based model introducing refined parameter estimation methods. *PLoS ONE*, 16(4 April).
- Watkins, C. J. C. H. (1989). *Learning from Delayed Rewards*. PhD thesis, King's College,



Oxford.

- White, L., Basurra, S., Gaber, M. M., Alsewari, A. A., Saeed, F., y Addanki, S. M. (2022). Agent-based simulations using genetic algorithm calibration: A children's services application. *IEEE Access*, 10:88386–88397.
- Wikstrom, K. y Nelson, H. T. (2022). Spatial validation of agent-based models. *Sustainability*, 14(24).
- Wilensky, U. (1999). Netlogo. <http://ccl.northwestern.edu/netlogo/>, Center for Connected Learning and Computer-Based Modeling.
- Wilensky, U. y Rand, W. (2015). *An introduction to agent-based modeling: modeling natural, social, and engineered complex systems with NetLogo*. MIT Press, Cambridge, MA., USA.
- Windrum, P., Fagiolo, G., y Moneta, A. (2007). Empirical validation of agent-based models: Alternatives and prospects. *Journal of Artificial Societies and Social Simulation*, 10(2):8.
- Witten, I., Frank, E., Hall, M., y Pal, C. (2017). *Data Mining: Practical Machine Learning Tools and Techniques*. The Morgan Kaufmann Series in Data Management Systems. Elsevier Science.
- Witten, I. H., Frank, E., y Hall, M. A. (2011). *Data Mining: Practical Machine Learning Tools and Techniques*. Elsevier.
- Wolf, I., Schröder, T., Neumann, J., y de Haan, G. (2015). Changing minds about electric cars: An empirically grounded agent-based modeling approach. *Technological Forecasting and Social Change*, 94:269–285.
- Wright, M. N. y Ziegler, A. (2017). ranger: A fast implementation of random forests for high dimensional data in C++ and R. *Journal of Statistical Software*, 77(1):1–17.
- Xiao, S. y Liu, R. (2021). Studies of covid-19 outbreak control using agent-based modeling. *Complex Systems*, 30(3):297–321.
- Xu, S., Chen, M., Feng, T., Zhan, L., Zhou, L., y Yu, G. (2021). Use ggbreak to effectively utilize plotting space to deal with large datasets and outliers. *Frontiers in Genetics*, 12.
- Xu, S., Chen, X., Xie, J., Rahman, S., Wang, J., Hui, H., y Chen, T. (2019). Agent-based modeling and simulation of the electricity market with residential demand response. *CSEE Journal of Power and Energy Systems*, PP(99).
- Xu, T., Gao, J., Coco, G., y Wang, S. (2020). Urban expansion in auckland, new zealand: a gis simulation via an intelligent self-adapting multiscale agent-based model. *International*

- Journal of Geographical Information Science*, 34(11):2136–2159.
- Yao, F., Zhu, J., Yu, J., Chen, C., y Chen, X. (2020). Hybrid operations of human driving vehicles and automated vehicles with data-driven agent-based simulation. *Transportation Research Part D: Transport and Environment*, 86.
- Ye, P., Chen, Y., Zhu, F., Lv, Y., Lu, W., y Wang, F. (2021). Bridging the micro and macro: Calibration of agent-based model using mean-field dynamics. *IEEE Transactions on Cybernetics*.
- Zhang, H., Vorobeychik, Y., Letchford, J., y Lakkaraju, K. (2016). Data-driven agent-based modeling, with application to rooftop solar adoption. *Autonomous Agents and Multi-Agent Systems*, 30(6):1023–1049.
- Zhang, J. y Sanderson, A. (2009). Jade: adaptive differential evolution with optional external archive. *IEEE Trans. Evol. Comput.*, 13:945–958.
- Zhang, Y., Grignard, A., Lyons, K., Aubuchon, A., y Larson, K. (2018). Real-time machine learning prediction of an agent-based model for urban decision-making (extended abstract). In *AAMAS*, volume 3, pages 2171–2173. International Foundation for Autonomous Agents and Multiagent Systems (IFAAMAS).
- Zhang, Y., Li, Z., y Zhang, Y. (2020). Validation and calibration of an agent-based model: A surrogate approach. *Discrete Dynamics in Nature and Society*, 2020.