



DARPA: SOCIALSIM AUDIT

September
2026

**DARPA can already simulate our clicks: how far
can it track our decisions?**

Paul JANIN & Jean LANGLOIS-BERTHELOT

SKEMA PUBLIKA

SKEMA Publika is SKEMA Business School's international think tank. It analyses major social, economic, technological and geopolitical changes to inform public and private decisionmaking. Drawing on the school's research and on scientifically validated external contributions, the think tank fuels public debate and issues recommendations for national and international decision-makers.

SKEMA Publika takes an interdisciplinary and international approach, enriched by SKEMA's global network of campuses and a community of experts from the academic and professional worlds. This international dimension is not simply a network, but a way of thinking. SKEMA Publika therefore connects local dynamics with global transformations to provide a decentralised and multipolar perspective on the major challenges of our time, rather than one centred on a single national lens.

The views expressed in this text are those of the authors alone.

© All rights reserved, SKEMA Business School, 2026

Cover: 'White concrete building under a cloudy sky during the day'. Harold Mendoza. Unsplash.com

SKEMA Publika

SKEMA Business School, Grand Paris Campus
5 Quai Marcel Dassault – PO Box 90067
92156 Suresnes Cedex, France

Tel.: +33 1 71 13 39 32
Email: publika@skema.edu
Website: www.publika.skema.edu

COLLECTION

‘Exploring the grey areas of geopolitics, digital technology and global risks’.

In the ‘Uncertainty’ collection, the works analyse tensions, emerging risks and dynamics of instability on a global scale

EDITORS

Frédérique Vidal is Director of Development at SKEMA Publika and Director of Strategy and Scientific Impact at SKEMA Business School. A full professor of Biology, she was president of the University of Nice Sophia Antipolis from 2012 to 2017, and subsequently served as minister of higher education, research and innovation in the Philippe and Castex governments from 2017 to 2022. She served as special advisor to the President of EFMD and is currently also the Permanent Representative of the Principality of Monaco to the United Nations Environment Programme and the Whaling Commission.

Sean Scull, Think Tank Project Manager, is a doctoral candidate in Information and Communication Science at Université Paul Valéry – Montpellier III. He holds a degree in Political Science with a specialisation in International Relations from the University of Gothenburg, and a Master’s degree in International Politics with a focus on Anglophone politics from the University of Toulon. Sean has lived and worked in Sweden and the United States

AUTHORS

Jean Langlois-Berthelot coordinates the 'Dual-Use Technologies and Innovations' module of the MTI Master's programme at INSTN-CEA/École des Mines, Paris Dauphine.

Paul Janin is a senior officer in the French Army, specialising in cognitive science, influence and cognitive warfare. As part of the Army Higher Military Scientific and Technical Education programme, he was tasked with analysing the scientific, academic and institutional ecosystem involved in these fields, carrying out fieldwork, analysis and exchanges with the sector's leading institutions. He notably spent time working at CEA Paris-Saclay, liaising with scientists involved in Defence Innovation Agency (AID) and Directorate General of Armaments (DGA) funding programmes relating to cognitive warfare and influence. He has published in several journals, including *Revue Défense Nationale*, *Polytechnique Insights* and *ISTE OpenScience*.

ABSTRACT

In this paper, Jean Langlois-Berthelot and Paul Janin present a case study focusing on DARPA's SocialSim programme. Drawing on a scientific and technical audit of its main publicly available components, they examine the associated publications, the code, the dependencies, the runtime conditions and the metrics employed. This analysis is supplemented by ten independent tests covering five cascade measures. The results confirm SocialSim's ability to represent digital events and their flow, whilst revealing the limitations of its publicly available components in terms of reproducibility. The audit thus establishes not only what SocialSim can effectively measure and simulate, but also what it cannot infer without additional evidence regarding human interpretation, causal influence and organisational decision-making.

TABLE OF CONTENTS

Abstract	i
Introduction.....	1
I. SocialSim: what DARPA really wanted to simulate	3
II. Methodology: how the audit was carried out	8
III. What the four repositories actually simulate.....	13
IV. What the audit reveals about the decision	16
Conclusion	20

INTRODUCTION

This paper is the second part of a two-part series on digital influence. The first part, entitled *'The Battle of Narratives'*, focused on the broader issues surrounding digital influence. This second paper aims to extend our analysis through a case study of DARPA and its *SocialSim* programme. DARPA (*Defence Advanced Research Projects Agency*¹), the advanced research agency of the US Department of Defence, funds exploratory programmes designed to overcome scientific or technical barriers. As part of its work, the agency completed *SocialSim (Social Simulation for Evaluating Online Messaging Campaigns*²) in 2021, with the aim of creating computer simulation tools capable of modelling the spread and evolution of information on social media. This programme addressed a specific question: is it possible to replicate and predict the activity of very large populations on digital platforms without reducing all users to the same behaviour? It should be noted that its objective focused on the spread and evolution of information online, not on the complete reconstruction of a society, human consciousness or a chain of command.

Thus, to trace the path from digital activity to real-world impact, we conducted a scientific audit of the four main public repositories (MACM, RHPC_SMPLE, GDELT_GA_SEARCH and socialsim_package³) derived from SocialSim. We examined the publications, code, dependencies, execution paths and metrics, and then verified several basic calculations separately. Our results confirm the value of SocialSim: the programme provides a particularly rich representation of digital events and their circulation. Human interpretation and the sequence of organisational decisions form a distinct layer, lying beyond its original scope. The study thus demonstrates how a science of online social dynamics can inform, without being conflated with, a science of decision-making.

¹ Home | DARPA. (2026). darpa.mil. <https://www.darpa.mil/>

² *Computational Simulation of Online Social Behaviour*. (2026). <https://www.darpa.mil/research/programs/computational-simulation-of-online-social-behavior>

³ UCF Social SIM. (2026). GitHub. <https://github.com/UCF-SocialSim>

In the first section, we will revisit the rationale behind SocialSim. In the second section, we will discuss the methodology used to carry out our audit of the programme. In the third section, we will examine what the four repositories actually simulate. In the fourth section, we will reveal what the audit demonstrates about decision-making.

I. SOCIALSIM: WHAT DARPA ACTUALLY SET OUT TO SIMULATE

The term ‘simulation’ here refers to an executable model, that is to say, a set of entities, states and rules that the computer evolves. In a multi-agent model, the entities are agents equipped with local rules: to post, reply, share, remain inactive or react to an external event. The collective phenomenon results from their interactions. A ‘cascade’ is the tree formed by an initial publication and the shares, responses or contributions that stem from it. Its size reflects the number of events; its depth measures the length of the relays; its width indicates how many branches appear at the same level; its duration tracks its development over time. These properties describe a flow. They do not convey the meaning that people attribute to the content.

SocialSim is not a single piece of software. Several teams have developed different models and tools as part of the programme. Our audit focuses on the four public repositories associated with Deep Agent, a project at the University of Central Florida. For the sake of clarity, we would like to point out that a repository is a version-controlled archive of code: it stores the files, their version history and, where applicable, the instructions needed to reproduce an experiment. The four repositories examined (MACM, RHPC_SMPLE, GDELT_GA_SEARCH and socialsim_package) cover four complementary functions: generating actions, representing platform-specific rules, searching for external events and comparing simulations with observed data.

The idea has become familiar: thanks to big data, artificial intelligence and multi-agent models, it is now supposedly possible to simulate a society, anticipate its reactions or test the effect of an intervention. It is an appealing prospect. It suggests a miniature world in which artificial agents receive information, form opinions, coordinate their actions and make decisions. All one would need to do is speed up time, alter a parameter and observe the resulting future.

A simulation necessarily involves selection. This implies that it is not a scaled-down copy of reality, but a construct that selects certain entities, certain states, certain relationships and certain transformation rules. Anything that does not fit this definition simply remains outside the scope of the experiment. Computing power then enables to explore the consequences of the chosen assumptions on a large scale; it does not, in itself, alter their scientific status.

The first task of scientific inquiry is therefore to ask what the model actually calculates. Does it count publications? Does it reproduce the structure of a network? Does it assign each agent a state called 'opinion'? Does it simulate exposure to content, its interpretation, an observable action or an organisational decision? These objects are not equivalent. A piece of software may excel at one, whilst the others fall within a different level of analysis.

Take, for example, a post shared a thousand times. The trace is simple: a thousand sharing events have been recorded. But these shares may express agreement, outrage, mockery, fact-checking, a desire to raise the alarm, the work of a bot or a coordinated response from an institution. The same visible action encompasses contradictory meanings. A cascade describes how content circulates; it does not automatically reveal what participants have understood from it or what they will do next.

The same distinction applies within an organisation. A message may be technically delivered without being noticed. It may be read without being understood, understood without being deemed credible, deemed credible without going through a validation process, validated without resulting in an order, or transformed into an order whilst certain reservations are set aside in the process. The volume of messages, their speed and their open rates therefore describe part of the decision-making process, but not the whole of it.

This distinction is crucial when simulation helps to guide action: crisis management, cybersecurity, intelligence, strategic communication, operations management, public policy or the management of a large organisation. A propagation map provides a precise answer to the question of circulation; it does not, on its own, answer the question of strategic

impact. A volume forecast can be highly effective without directly measuring a change in behaviour. A psychologically realistic interface, too, retains the status of the assumptions that define the agents' internal states.

The difficulty does not stem solely from the software. It stems from our tendency to let words slip. A binary variable is termed a 'belief'; crossing a threshold becomes 'persuasion'; sharing becomes 'adherence'; a drop in traffic becomes an 'effect'; a convergence of behaviours becomes 'coherence'. The name of the variable ends up serving as proof in itself. Yet the same threshold- mechanism could just as easily represent the adoption of a tool, the activation of an alarm or participation in an activity. External evidence is required to determine which interpretation is justified.

The strength of a social simulation thus depends on the chain of evidence linking its variables to its conclusions. The most robust models make visible what they represent, separate observation from inference, link each output to a specific assertion, and spell out their scope of validity. SocialSim offers precisely such a body of publicly available data, of a rarity in its scope, to examine this chain without reducing the programme's value to a general assessment.

This ambition addressed a real challenge. Top-down models can easily describe millions of entities based on global laws, but often oversimplify the differences between individuals. Bottom-up models assign distinct rules to agents and allow collective phenomena to emerge from their interactions, but quickly become computationally expensive. SocialSim sought to combine the two advantages: behavioural heterogeneity and scalability.

The challenge conducted on GitHub, a platform where developers host code and log their changes, illustrates this level of ambition. The teams had thirty months' worth of metadata, dates, actors, types of action and objects involved, and had to forecast activity for the following months in an environment comprising around three million users, thirty million actions and six million deposits. Several families of agents were compared against metrics relating to volumes, actions and their targets. The most useful finding is also the most instructive: the most complex model was not consistently the best. A stable statistical profile specific to each user often

produced good predictions, particularly because behaviour on GitHub evolved relatively slowly⁴.

This result demonstrates what good computational science can achieve. Models are compared against future data; complexity is not rewarded as a matter of principle; a simple baseline can outperform a sophisticated architecture. It also highlights what experience does not prove. Predicting that a user will open an issue, edit a repository or reply to a conversation does not mean that we know their beliefs, intentions or the offline effects of their action.

We have fixed the scope of our audit as of 17 July 2026. It comprises MACM, RHPC_SMPLE, GDELT_GA_SEARCH and socialsim_package. Together, these repositories form less an integrated product than a distributed laboratory. Each addresses a different question.

Public repository	Main function	Established scientific contribution	Identified scope limitations
MACM	Generating the next steps in a conversation	Distinction between network pressure, external shock and spontaneous activity	Observable behaviour alone does not determine human interpretation
RHPC_SMPLE	Representing platforms, their actions and their visibility flows	Platform-specific characteristics of actions, content and flows	The digital environment, on its own, does not reconstruct an organisation's internal communications
GDELT_GA_SEARCH	Search for external events correlated with social activity	Distinction between internal dynamics and external context	Temporal correlation alone does not establish a causal effect

⁴ Blythe, J. et al. (2019). 'Massive Multi-Agent Data-Driven Simulations of the GitHub Ecosystem'. In *Advances in Practical Applications of Survivable Agents and Multi-Agent Systems*, 3–15. https://doi.org/10.1007/978-3-030-24209-1_1

Public repository	Main function	Established scientific contribution	Identified scope limitations
socialsim_package	Comparing simulations and observations using metrics	Characterisation of volumes, cascades, networks, groups and temporalities	Traffic metrics do not directly measure human impact

The audited scope is significant: four repositories, three technical ecosystems, at least 39,800 documented lines excluding unmerged Java components, 595 versioned changes and forty localised observations. RHPC_SMPLE identifies six platforms, including Twitter, Reddit, YouTube, Telegram, GitHub and Facebook. No comprehensive automated test suite was identified within the scope of the four repositories. These figures describe the scope of the audit; they do not constitute a quality rating. Three *commits* (an IT term referring to the recording of source code⁵) may contain a correct mechanism, whilst four hundred may still contain an error. An observation may indicate a technical incompatibility, a path likely to fail or a validation constraint; it is not automatically a ‘bug’ that has affected a published result.

The audit reveals a robust intellectual architecture and precise scope limitations. MACM distinguishes between several possible causes of a single action. An activity may result from endogenous pressure within the network, an exogenous shock or spontaneous internal activity. This breakdown prevents any increase in activity from being mistaken for a contagion phenomenon. RHPC_SMPLE treats the platform as an active environment: the available actions, visibility flows, content types and timeframes differ across digital spaces. GDELT_GA_SEARCH utilises GDELT, a global database that indexes events and press articles, to investigate whether online activity accompanies an external event. Finally, SocialSim_package provides a rich set of metrics for comparing simulated and observed cascades. None of these observations detracts from the programme’s contribution: they simply indicate the extent to which each result allows for interpretation.

This is the scientific core of SocialSim: explicit distinctions, event formats, generative mechanisms and comparison scenarios. An examination of the

⁵ Neodelta. (2026). Commit – definition | Neodelta. <https://neodelta.eu/glossaire/commit>

code helps to clarify the conditions under which these findings can be replicated, as well as the boundary beyond which the circulation of a signal is no longer sufficient to establish its interpretation or effect. This boundary does not represent a failure of the programme; it corresponds to the difference between its specific purpose and the broader conclusions one might be tempted to draw from it.

II. METHODOLOGY: HOW THE AUDIT WAS CONDUCTED

We carried out the audit from the compilation of the corpus through to the inspection of mechanisms and the execution of independent tests. This is neither a general assessment of DARPA nor a ranking of software. Our method brings together four elements that are often examined separately: what the publications claim, what the code actually represents, what a specific experiment manages to achieve, and what the outputs actually allow us to conclude.

The first challenge was to establish a common language. A platform records clicks, posts and links; an organisation produces alerts, assessments, judgements and orders. Using the terms 'information', 'influence', 'behaviour' or 'decision' interchangeably would have rendered any comparison misleading. We therefore organised the observations into six layers, ranging from a simple recorded trace to the reproduction of decisions. This framework does not add any phenomena to the code: it indicates what each level establishes and the exact limits of the evidence it provides.

The scope of SocialSim, and, by extension, that of many so-called social intelligence solutions, is broken down into six layers. These do not constitute six separate worlds, but six nested levels of evidence. Each layer builds on the previous one and provides different information.

Layer	What is actually observed	Limit of the evidence at this level
1. Trace	A click, a message, a log, a written order or a timestamp exists	The existence of the event does not establish either its understanding or its use

Layer	What is actually observed	Limit of the evidence at this level
2. Circulation	Content has travelled at a certain speed, via a particular network or through a chain	Dissemination does not imply a change in interpretation or decision
3. Reception	An actor has been exposed to, opened or processed a signal	Exposure does not result in a change to one of the organisation's premises
4. Interpretation	A situation is assigned a classification and a degree of confidence	A local classification does not imply its adoption by the organisation
5. Decision-making communication	An alert, a hypothesis, a decision or an order becomes available for further action	An isolated communication does not ensure the consistency of the entire chain
6. Decision-making replication	A decision becomes the premise for further communications, corrections or closures	The observed trajectory does not, on its own, establish causal attribution

The first layer is the trace. It is essential, because without it the investigation relies on memories or accounts that are impossible to verify. But a trace does not contain its own meaning. A log proves that an operation took place within a system; it does not prove that an attack occurred. An order recorded in an incident log proves that a statement was recorded; it proves neither that it was understood nor that it was carried out.

The second layer is circulation. Volumes, centralities, the importance of an actor based on their position in the network, cascades and speeds describe the trajectory of a signal. SocialSim is particularly strong in this regard. It can compare the size of a cascade, its depth, its breadth, its duration or its structural virality. Two cascades of the same size are not alike if the first radiates out from a central source whilst the second progresses via successive relays. This difference is scientifically informative. It still does not indicate whether the relays endorse the content or contest it.

The third layer is reception. In a computer system, it can be established that an event has been delivered to a simulated agent, stored in its memory or made visible in its feed. In humans, availability precedes attention, understanding and memorisation. A probability of attention represents this

transition in the model; its human significance is reinforced when it is cross-referenced with independent observations.

The fourth layer is interpretation. An individual or group assigns a form to the situation: a breakdown or an attack, a local incident or a strategic threat, confirmed information or a tenuous hypothesis. This layer is not merely a hidden mental state. It may leave traces in a note, a classification, an alert level or a request for verification. However, a local interpretation does not automatically become an organisational interpretation. Its adoption through a communication that alters the possibilities for the next steps constitutes precisely the organisational threshold.

The fifth layer is therefore decision-making communication. An alert singles out a phenomenon. A hypothesis offers an explanation. A challenge maintains a divergence. A ruling temporarily rules out certain options. An order authorises an action. A decision to remain silent may itself constitute organised communication. The focus is no longer on the number of messages, but on the way in which a selection becomes available, transmissible and binding for other functions.

The sixth layer is decision-making reproduction. An organisation is not coherent simply because it produces a single final decision. It is coherent when its successive decisions can be linked to one another, retain useful uncertainties and reopen a distinction that has become insufficient. The first decision provisionally closes off possibilities; it then becomes a premise for the second. The second decision may confirm, interpret, correct or abandon the first. Coherence lies in the quality of this chain.

This conception changes the very definition of decision-making. Deciding does not merely consist of choosing from among pre-existing options. The organisation first transforms an indeterminate situation into a manageable problem. It selects signals, distinguishes between several hypotheses, assesses their credibility, sets a threshold for action and sometimes constructs the options it will compare. John Dewey had already described decision-making as an inquiry that transforms a problematic situation⁶. Herbert Simon and James March subsequently demonstrated the extent to

⁶ Dewey, J. (1993). *Logic: The Theory of Inquiry*. PUF.

which this inquiry depends on the limits of attention, time, routines and procedures⁷. Niklas Luhmann's organisational theory adds a crucial insight: the organisation reproduces itself through decision-making communications that become the premises for subsequent communications⁸.

Decision-making coherence therefore refers neither to unanimity nor to speed. An organisation can remain coherent whilst maintaining explicit disagreement, provided that this disagreement is localised, communicated and linked to conditions for revision. Conversely, total agreement can become fragile if it stems from premature closure or the suppression of reservations. A rapid response can be excellent if it remains reversible; it becomes riskier when it makes public an assessment that is still uncertain and reduces the scope for correction.

Two chains of reasoning can lead to the same outcome whilst possessing opposing qualities. In the first, several assumptions are retained, the degree of confidence is explicit and the chosen course of action can be corrected. In the second, a possibility becomes a certainty without new information, reservations disappear from the records and the action results in an irreversible commitment. The final outcome is identical. Vulnerability to the next disruption, however, is not.

Following this separation of the six layers, the second part of the audit traces the chain leading from an idea to a conclusion: conceptual model, implementation, experiment and interpreted result. This method avoids attributing to the programme as a whole a property that is merely stated in an article, located in a file or obtained during a particular execution.

The conceptual model defines entities, states and mechanisms. The implementation translates these choices into code and data. The experiment associates a specific version of the code with parameters, initial conditions, an environment and random seeds. Finally, the interpreted result transforms an output into a statement. An error may occur at any stage. A publication sets out the scientific intent but often omits the details of execution. A repository shows part of the implementation but may

⁷ March, J and Herbert, S. (1993). *Organisations*. John Wiley & Sons.

⁸ Luhmann, N. (2000). *Organisation and Decision-Making*. Westdeutcher Verlag.

contain multiple versions, abandoned scripts and functions that were never called. The fact that a programme runs does not prove that its model corresponds to reality. The fact that a result resembles the data does not prove that the correct mechanism produced it.

The classic distinction between verification and validation remains fundamental here. Verification involves asking whether the model has been constructed correctly: does the code implement the stated rule? Validation involves asking whether the model has been constructed for its intended purpose: do the simulated entities and behaviours correspond sufficiently to the phenomenon under study? A model may be valid for comparing two stylised mechanisms, but invalid for predicting a real-world crisis. It may be excellent for teaching but unsuitable for operational decision-making.

For this reason, the protocol records the repository, commit or version, capture date, environment, dependencies, the command executed and the data used. The name of a piece of software is not sufficient evidence in itself. An open licence provides information on the legal possibilities for reuse, not on scientific validity. A programme's reputation does not guarantee the reconstructibility of its artefacts; the age of a piece of code does not render its concept false. The audit thus distinguishes between the scope of the mechanism, the conformity of its implementation and the actual conditions of its execution.

Each observation is assigned a level of evidence. Level zero means that a piece of information has not been found in the sources inspected; it does not allow the conclusion that it does not exist. Level one corresponds to a statement in an article or documentation. Level two identifies the mechanism in the code or a configuration. Level three adds a recorded execution result. Level four compares the result with an independent reference: a manual calculation, an analytical solution, a second implementation or separate calibration data. These levels do not form an overall score. They merely indicate the extent to which a claim has been verified.

The absence of an overall score is a methodological choice. Adding up the number of tests, the licence, maintenance, behavioural richness, speed and empirical accuracy would produce a dimensionless figure. A system may be an excellent analytical testbed, a poor software foundation and a

remarkable educational tool. Another might be very well maintained but contain no domain theory. The matrix preserves these profiles rather than averaging them out into an arbitrary figure.

The final axis of the audit classifies elements according to four statuses: retainable as a definition or metric; re-implementable from an independent specification; isolatable as a technical dependency; not selected for direct integration. This classification is not a quality ranking. It distinguishes the scientific value of an idea, the state of its historical code and its suitability for another object. An engine may thus be unsuitable for direct integration whilst remaining intellectually fruitful; a metric may be clear whilst its implementation requires adaptation; a dependency may be technically functional whilst supporting a different level of inference.

Each observation thus links a localised finding to a risk of misinterpretation. When the same variable represents both exposure and adoption, the problem lies not only in the code: the output becomes unable to distinguish between content viewed and behaviour adopted. When a random number generator is uncontrolled, the exact repetition of the experiment becomes indeterminate. When a dependency makes undocumented network calls, the actual scope of execution is not captured in the manifest. The audit does not transform these findings into a general verdict; it demonstrates the precise consequence of each one based on the available evidence.

III. WHAT THE FOUR REPOSITORIES ACTUALLY SIMULATE

Technical inspection complements scientific analysis by tracing the path linking a hypothesis to a result. A dependency is an external library whose code is required for the application to function; when it changes or disappears, the experiment may become difficult to reproduce. A runtime environment brings together the operating system, the programming language, the libraries and, in some cases, the required hardware. An automated test verifies that a specific property remains true following a modification. When a complete test suite is not available, verification is carried out using independent and explicitly calculated test cases.

The four repositories draw on three technical domains. MACM uses Python, a common language in data science, and CUDA, an NVIDIA technology that performs calculations on a graphics card. RHPC_SMPLE is based on C++, MPI, a standard enabling multiple computers or processors to communicate during a calculation, and Repast HPC, a distributed multi-agent simulation framework. GDELT_GA_SEARCH is written in Java. Socialsim_package uses Python and libraries such as NumPy and Pandas to manipulate data tables. These details are important because a model may remain scientifically interesting even when its historical context has become difficult to reconstruct.

The overall result highlights four complementary achievements: SocialSim represents heterogeneous actions, simulates on a large scale, integrates multiple sources of dynamics, and compares its outputs with future data. The public artefacts examined were designed for this specific purpose. Extending them to reflect the coherence of an organisation would require a different level of representation, which was not the focus of the programme.

MACM offers the most fruitful concept: a visible action can be triggered by network pressure, an external shock or spontaneous activity. This decomposition avoids automatically attributing all activity to propagation. Its public implementation reflects an outdated research environment: partial dependencies, a minimal scenario reliant on external data, direct selection of a CUDA device, Pandas operations that have become obsolete, and the interpretation of certain text inputs by ``eval``, a function that executes text as code. The review also identified several paths warranting isolated verification: a shock variable apparently confused with that of endogenous influence, an empty loop using default values, and a possible out-of-bounds access during memory mixing. These observations do not call into question the published scientific results; they delineate what can be re-executed directly from the public repository. The breakdown between network, shock and internal activity remains clearly identifiable and constitutes a key contribution.

RHPC_SMPLE provides the most comprehensive description of platforms as distinct environments. It separates users, content, actions, selection strategies and visibility flows. This distinction is crucial: posting on Twitter, replying on Reddit, editing on GitHub or commenting on YouTube does not create the same trace or the same temporality. The C++ engine is based on a distributed computing stack combining MPI, Repast HPC, NetCDF and

Boost. Its build image is based on Ubuntu 16.04 and its public documentation leaves certain sections open; reconstruction therefore requires recreating a specific historical environment. Certain rules governing selection and memory management also require isolated testing. The conceptual contribution remains clear: the platform is not merely a graph, but an environment that shapes what participants see and are able to do.

GDELT_GA_SEARCH provides an essential methodological reminder: a surge in activity may be linked to an external event rather than to internal contagion. The repository uses a genetic algorithm to search for GDELT queries whose time series correlate with social activity. As with any exploration of a large number of queries, this method carries a risk of overfitting: the pattern that best fits the past is not necessarily the one that will persist thereafter. The use of a random generator without a controlled seed, absolute local paths and certain peculiarities of the temporal processing make exact reproduction more challenging. The repository firmly establishes the relevance of an exogenous context; causal attribution is the subject of a complementary protocol.

Socialsim_package is the component most directly reusable for scientific purposes. It calculates the size, depth, breadth, duration and virality of cascades, as well as activity concentration, network components, modularity – the degree to which a network divides into subgroups –, persistent groups, bursts and various metrics comparing simulation and reality. Its historical Python environment requires some adaptations to accommodate recent versions of the libraries. The aim of the audit is therefore to verify that the modern port retains the original definitions exactly.

We therefore wrote an independent core limited to five basic metrics: size, depth, maximum width, duration and structural virality. Structural virality does not measure the popularity of content; it describes the pattern of its propagation by calculating the average distance between the nodes in a cascade. Before any interpretation, we calculated the expected results for simple cases: an empty cascade, a single node, a chain of four nodes, a five-node star and a balanced binary tree with seven nodes. Five further tests verify the order of inputs and the rejection of invalid structures. All ten tests were successful. This result is deliberately limited in scope. It proves that five calculations are correct within a defined scope. It does not validate an agent's psychology, the effect of information, or a decision-making architecture.

The deliberately narrow scope of this test makes it an intelligible proof. A historical implementation and an independent implementation can be compared on the same micro-cascades; any divergence can then be pinpointed. However, this result does not extend to MACM mechanisms, the states assigned to agents, or the effects of an external shock. It verifies a family of calculations, not a theory of behaviour. This distinction between local analytical proof and general validity constitutes one of the audit's most significant findings.

The limitations observed can thus be divided into three categories. The first concerns the immediate reproducibility of public repositories: legacy dependencies, inconsistent documentation, historical environments and partial test coverage. The second concerns the scope of the verification: the ten independent tests validate five cascade metrics, not the full set of behavioural mechanisms. The third relates to the scope of inference: traces and correlations powerfully describe digital activity, but do not, on their own, establish human interpretation, the causality of an intervention or collective decision-making. These limitations are real, but they do not constitute a case against SocialSim. The first limitation stems largely from the nature of archived research code; the other two relate to standard rules of evidence and the very scope of the programme. SocialSim remains a leading contribution to the modelling of observable events and their circulation, with a rare ambition for large-scale comparison. Its specifications, metrics, causal distinctions and reference cases retain lasting significance.

IV. WHAT THE AUDIT REVEALS ABOUT DECISION-MAKING

A further distinction clarifies operational uses: predicting, explaining and measuring the effect of an intervention are three distinct operations. A model may correctly predict tomorrow's volume of activity because that volume resembles yesterday's. A different method is required when estimating what would happen if an organisation were to issue a denial, suspend a service or choose not to respond.

The concept of ‘effect’ implies a comparison between the observed action and a credible alternative. What would have happened without a statement? With a statement issued later? With a different wording? A drop in volume following a communication does not prove that the communication has calmed the conversation: the momentum may already have been waning. Nor does an increase prove failure: speaking out may have focused the controversy whilst reducing dangerous behaviour. Raw visibility alone is therefore not sufficient to assess the outcome.

Randomised trials provide the strongest comparison where possible. The study by Bond and his co-authors, involving 61 million Facebook users, was able to distinguish between a direct effect and social transmission because exposure was based on experimental allocation⁹. In most real-world crises, such allocation is impossible. The analysis therefore relies on cautious comparisons: similar units, staggered time windows, synthetic controls, simulated scenarios or process variants. These methods seek a counterfactual, that is to say, a credible estimate of what would have happened in the absence of the action under study. None of them creates a perfect parallel world; each merely makes certain interpretations more or less tenable.

SocialSim can generate multiple scenarios, whilst its public repositories focus on events, volumes, participants and forms of cascade. A specific organisational intervention, its functional author, timing, justification, alternatives and target behaviour, falls under a complementary causal protocol. This difference reflects SocialSim’s specialisation in online social dynamics, not a shortcoming in relation to its original objective.

The effect is particularly evident in the chain of events preceding the final action. A disruption may not immediately alter this action, but it may shift an emergency threshold, erode trust in a channel, reduce the number of retained hypotheses, or cause a public communication to be issued prematurely. Conversely, a change in behaviour may stem from a logistical constraint without any shift in interpretation. The chain of premises therefore provides a more precise result than the final action.

⁹ Bond, R. M., Fariss, C. J., Jones, J. J., Kramer, A. D. I., Marlow, C., Settle, J. E., & Fowler, J. H. (2012). A 61-million-person experiment in social influence and political mobilisation. *Nature*, 489(7415), 295–298. <https://doi.org/10.1038/nature11421>

This approach allows us to study saturation without resorting to naive arithmetic. A very high volume of converging signals may remain manageable. A few contradictory messages between critical functions can disrupt decision-making. The problem is not merely the quantity of information, but the cost of translation, the divergence of qualifications and the ability to maintain explicit uncertainty. An organisation is not saturated the moment it receives 'too many' messages according to a universal threshold; it is saturated when its communications can no longer produce compatible distinctions within the time available.

It also helps us understand that restoring coherence does not always require more information. The organisation can recover by correcting a map, changing the medium, slowing down a process, re-establishing a distinction or revisiting a previously discarded hypothesis. The right course of action may be clarification, translation or organised silence. Adding data to a system that has lost its categories can exacerbate the confusion.

Finally, this perspective reveals the specific effects of measurement. A unit assessed solely on its speed may eliminate uncertainty in order to close cases more quickly. If assessed on the number of alerts processed, it may fragment events or close them prematurely. If assessed on convergence, it may eliminate useful disagreements. The indicator then becomes a disruption to the very system it purports to observe. Coherence no longer presents itself as a sovereign variable, but as a reasoned interpretation based on dimensions that may diverge.

Comparing SocialSim with decision-making coherence ultimately reveals four theoretical findings: saturation is not merely a matter of volume; uncertainty follows a trajectory; restoring coherence does not always require more information; and the effect of a disturbance can be seen in the transformation of the premises. These findings do not describe a future architecture; they clarify what becomes observable when we cease to equate activity, influence and decision-making.

The first finding concerns saturation. Circulation models naturally represent volumes, rhythms, densities and cascading . They thus favour a quantitative definition of overload: the system would be saturated beyond a certain number of messages or events. An analysis of organisational

communications leads to a different conclusion. Volume only becomes critical when it prevents the production of compatible distinctions or when the cost of translation absorbs available capacity. An intense but convergent flow may remain manageable; a few contradictory signals between central functions may be enough to destabilise the whole system. Density and divergence therefore do not describe the same phenomenon.

This distinction sheds light on the specific scope of SocialSim. The data sets provide a detailed representation of the number of events, their relationships and their distribution over time. The instance in which two units employ incompatible classifications, as well as the work required to convert a technical diagnosis into a legal or operational category, belong to an additional organisational layer. SocialSim makes the saturation of activity visible; decision-making analysis extends this result by observing the transformation of messages into premises.

The second finding concerns uncertainty. In a conventional analysis, it often appears as a property attached to a state: the probability of an event, confidence in a prediction, or the spread of a distribution. Within an organisation, it also follows a trajectory. A hypothesis may be passed on with its degree of confidence, reformulated without it, reinforced by an authority figure without any new information, or converted into certainty at the time of a debriefing. The loss of consistency is therefore not an initial calculation error; it occurs during the internal circulation of the assessment.

The difference is evident in a simple case. A technical unit estimates the probability of an attack at 40 per cent and recommends verification. The next level up records it as 'possible attack'. The subsequent briefing mentions 'probable attack'. The final order treats the attack as an established fact. Each statement may appear consistent with the previous one, whilst uncertainty has been gradually erased without any new observations. Records of volume and dissemination do not capture this semantic drift. An organisation may therefore have all the initial data at its disposal and yet lose the decisive piece of information: the fragility of the assessment.

The third finding concerns the restoration of coherence. The decision-making mindset spontaneously associates improvement with the addition of information: more sensors, messages, calculations or expertise would

produce a better representation. Yet an organisation can regain a coherent trajectory- without receiving any new data. A correction to a map, a change of medium, the reopening of a hypothesis, the re-establishment of a distinction or the temporary slowing down of a chain may suffice. Restoration is therefore a reconfiguration of the relationships between communications, not an increase in their volume.

This explains why silence is not always a void and why disagreement is not always a breakdown. Organised silence can prevent a fragile qualification from becoming irreversible. Explicit disagreement can preserve two hypotheses until a distinguishing factor emerges. Conversely, an increase in the number of messages can accelerate the premature closure of a problem. The usual indicators of activity, speed and convergence may thus appear to be moving in a seemingly favourable direction, whilst actual reversibility is decreasing.

The fourth finding shifts the measurement of the effect towards the chain of premises. An external disruption does not necessarily alter an individual's belief in order to influence a collective decision. It may change what is deemed urgent, credible, publishable or actionable. It may alter the trust placed in a channel, shift an escalation threshold or impose a visibility metric as a definition of severity. The organisation then continues to communicate and act, but on the basis of transformed distinctions.

This framework allows us to compare two sequences that result in the same behaviour. In the first, a tentative alert is retained as a hypothesis, several options remain open and the chosen course of action is reversible. In the second, the same alert becomes a certainty, closes off alternatives and leads to a public commitment. If the analysis stops at the final action, the two sequences are identical. If it traces the premises, they present opposing structures. The effect no longer lies solely in what has been done, but in the possibilities that each communication has left open for the next.

CONCLUSION

These four findings explain SocialSim's precise position within the field. The programme observes and generates external disturbances with great richness: actions, networks, temporalities, platforms, groups and cascades.

It can represent the technical availability of content and infer certain latent states, that is to say, properties that are not directly observed but estimated on the basis of actions. Qualifications, the degrees of trust conveyed, trade-offs, conditions for revision and premise effects fall within a second scientific object of study: the reproduction of organisational decisions. The two objects are complementary; the second begins where the first has already provided a structured representation of the digital environment.

This distinction also clarifies the status of the seven dimensions introduced by the analysis. Selection denotes what becomes relevant; uncertainty tracks the evolution of the degree of confidence; translation observes shifts in meaning between functions; coupling describes the compatibility of qualifications; temporality concerns response times; reversibility indicates possibilities for correction; the premise effect shows what a communication subsequently makes possible or impossible. None of these dimensions is to be confused with a mere volume of messages. They describe the intermediate transformations that are missing between a recorded disruption and a final decision.

Nor can they be summarised by a single index. A chain may be rapid but irreversible, coordinated but based on an erroneous assessment, or cautious but incapable of producing a follow-up. The dimensions may diverge without one automatically providing the deciding factor. This absence of a definitive score is not a flaw in the formalisation. It corresponds to the very structure of the object: decision-making coherence is a property of a trajectory, not a simple quantity attached to a single moment in time.

These four findings allow us to conclude the analysis of SocialSim's precise contribution. The DARPA programme has given a rigorous computational form to a major part of the digital world: events, platforms, networks, cascades, groups and external signals. It has shown that competing models can be evaluated against future data and that sophistication does not always outperform a simple benchmark. This scientific legacy is substantial and clearly identifiable.

The audit highlights the continuity across several levels. SocialSim natively represents traces and flows, and enables a technical interpretation to be constructed. Its latent variables offer hypotheses regarding human

interpretation, which external data can then consolidate. Decision-making communication and the reproduction of an organisation's decisions extend this architecture to another level. Accurate activity forecasting, cascade measurement, inference about states, and decision analysis are therefore not competing outcomes, but successive links in the same chain of inquiry.

The four repositories make different contributions. Event formats, the plurality of platforms, cascade metrics and the distinction between internal dynamics and external shocks have clear scientific significance. The historical engines document the technical conditions under which these ideas were implemented. The audit thus distinguishes between what constitutes a definition, a verified calculation, a localised mechanism, a bounded implementation or a still-hypothetical interpretation, without reducing these dimensions to a single rating.

Decision-making coherence is thus revealed for what it is: not perfect agreement amongst all stakeholders, but a system's ability to temporarily reduce complexity without undermining the conditions for its own revision. A coherent organisation does not speak with one voice. Its communications highlight its differences of opinion, preserve what remains uncertain, justify the action taken and maintain the possibility of revisiting a distinction that has become false.

The phrase in the title can therefore be taken seriously. SocialSim is already capable of counting, linking and generating a considerable proportion of our social traces. This achievement provides the foundation from which analysis can track the transformation of a situation into a problem, a problem into a decision, and an error into a correction. SocialSim is therefore not an incomplete model of an organisational decision; it is a fully-fledged model of the digital dynamics upon which that decision may need to act.

skema
THINK TANK

PUBLIKA

SKEMA Publika

SKEMA Business School, Grand Paris Campus
5 Quai Marcel Dassault – PO Box 90067
92156 Suresnes Cedex, France

Tel.: +33 1 71 13 39 32
Email: publika@skema.edu
Website: www.publika.skema.edu