

BUILDING A BETTER FOUNDATION FOR AGI

WHITEPAPER

2025

Authors:

Dr. Mohamed Seddik

Lead Researcher, AI and Digital Science Research Center
Technology Innovation Institute

Dr. Hakim Hacid

Chief Researcher, AI and Digital Science Research Center
Technology Innovation Institute

Introduction

The pursuit of artificial general intelligence (AGI) capable of advanced thinking and reasoning has been an elusive goal since the early days of AI. For almost sixty years, researchers tried multiple approaches that later failed to scale for various reasons.

ChatGPT's incredible success in generating authoritative responses in late 2022 reinvigorated interest in AGI research. In recent years, new generative AI techniques built on large language models (LLMs) have sparked enormous excitement about a new path to AGI. New models such as GPT-4 and open models such as the Technology Innovation Institute's (TII) Falcon series have demonstrated unprecedented abilities in understanding and generating human-like text. Many experts believed that AGI could soon be achieved by simply scaling up the data sets and computing resources for training LLMs. This approach has seen some progress but has also suffered from diminishing returns.¹

Now, researchers are exploring other approaches for teaching models to reason and use external tools. TII is actively engaged in foundational research that will accelerate the path toward the eventual development of AGI. Some of the many building blocks include new foundation models, reasoning models, reinforcement learning techniques for training reasoning models, creation of verifiable data, building verifier models, and support for tool use. This work is strategically aimed at creating the necessary foundations for AGI.

One promising approach has been the development of new reasoning models built on top of foundation models. The core idea is to separate the knowledge base and reasoning ability learned in a reasoning model from the active reasoning processes required to complete a complex task. Layering reasoning capabilities on a solid foundation will move us toward general problem-solving without an exorbitantly larger base model. This white paper explores how we can build a better foundation for AGI by focusing on refining reasoning approaches outside of traditional brute-force scaling techniques.

Key aspects include 1) new training approaches for reasoning skills; 2) test-time scaling for thinking token generation; 3) improving reasoning path search to explore multiple reasoning paths during inference; 4) curating verifiable data, and 5) building the infrastructure to support tool use for LLMs. Each of these components addresses a facet of reasoning challenges. Together, they form a blueprint for moving from today's LLMs toward more robust, reasoning-enabled AGI.

[1] M. Zeff, "Current AI scaling laws are showing diminishing returns, forcing AI labs to change course," TechCrunch. Accessed: Apr. 08, 2025. [Online]. Available: <https://techcrunch.com/2024/11/20/ai-scaling-laws-are-showing-diminishing-returns-forcing-ai-labs-to-change-course/>

Foundation model building blocks

A strong foundation model creates bedrock to support more advanced reasoning capabilities. Large pre-trained foundation models learn linguistic knowledge and patterns. Early models could serve as chatbots to answer questions or summarize text. Now, researchers are exploring how to train and fine-tune models that can also support reasoning processes. Initial research indicated that this required models with hundreds of billions of parameters.²

The next frontier lies in discovering ways to shrink models capable of supporting reasoning tasks to a much smaller scale. For example, TII's current efforts focus on developing high-quality, large-scale language models (LLMs) known as the "Falcon models." Early versions focused on instruction-following versions that could function as chatbots.

The latest released series is Falcon 3, which encompasses a range of model sizes from 1 billion to 10 billion parameters that can also enhance reasoning processes.³ This includes base and post-trained variants for instruction following use cases. The development process includes training base models and continuously improving their capabilities. These robust base models are crucial for more advanced applications, particularly those involving reasoning and reinforcement learning techniques.

Despite their small size, the Falcon 3 family was optimized for good performance in basic reasoning, coding, math, and understanding tasks.⁴ These models are also multi-lingual, and TII plans to extend Falcon 3 into multi-modal domains to support images, audio, and video. By open-sourcing performant models at various scales, the TII provides the AI community with flexible foundation building blocks to experiment with reasoning techniques on top. In the future, the TII will continue to advance foundation models, such as extending model size and exploring different architecture designs (pure-transformer, mixture of experts, hybrid transformer-mamba architecture), and add support for multi-modal understanding. These essential building blocks set the stage for better reasoning systems.

Most foundation model research focuses on the transformer architecture, first discovered by Google researchers in 2017 and widely used in most commercial large language models today. However, exploring the potential for other architecture designs that might be faster, more accurate, or perform better for particular use cases is essential.

[2] J. Wei et al., "Chain-of-Thought Prompting Elicits Reasoning in Large Language Models," Jan. 10, 2023, arXiv: arXiv:2201.11903. doi: 10.48550/arXiv.2201.11903.

[3] "Falcon 3: UAE's Technology Innovation Institute Launches World's most Powerful Small AI Models that can also be run on Light Infrastructures, including Laptops | Technology Innovation Institute." Accessed: Apr. 08, 2025. [Online].

Available: <https://www.tii.ae/news/falcon-3-uae-technology-innovation-institute-launches-worlds-most-powerful-small-ai-models>

[4] 12/17/2024 Team Intel, "Intel AI Solutions Support Falcon 3 Models," Intel. Accessed: Apr. 08, 2025. [Online].

Available: <https://www.intel.com/content/www/us/en/developer/articles/technical/intel-ai-solutions-support-falcon-3-fdn-models.html>

Foundation model building blocks

For example, as part of its work on the Mamba architecture, TII is exploring using state space models and hybrid designs, potential speedups, and further scaling of the context window, which is particularly essential for test-time scaling (i.e., generation of reasoning tokens). TII's Falcon Mamba 7B is a 7-billion parameter model using a sparse SSM architecture instead of a transformer architecture.⁵

Transformers are built to process correlations between all tokens, which yields a quadratic complexity in terms of the context size. With the state space approach, Mamba allows context representation in vector space, yielding linear complexity in terms of the context size, more efficiently for faster computing, and requiring less memory.

Although this approach is not fully competitive with transformer models for all tasks, it does show promise. Moreover, the research team at TII is working on building larger transformer-based models while exploring hybrid designs that combine both the advantages of the transformer and SSM designs, allowing a better trade-off between performance and inference speed, and other ongoing efforts focus on enabling the mixture of expert designs for improving efficiency.

At the same time, TII has continued to drive progress in scaling conventional LLMs while maintaining an open approach by releasing open-weight versions of the Falcon models to the research and development communities. While the training datasets are not entirely open-sourced, a portion was released to enable the broader research community to build upon its work. This fosters collaboration and improves external contributions and feedback.

[5] "UAE's TII Revolutionizes AI with New Falcon Mamba 7B Model." Accessed: Apr. 08, 2025. [Online]. Available: <https://www.tii.ae/news/uaes-technology-innovation-institute-revolutionizes-ai-language-models-new-architecture>

Creating a foundation for reasoning

With solid base models in place, the next step is enabling reasoning on top of them. This includes getting the models to generate and evaluate intermediate logical steps rather than just jumping from question to answer. This can be enabled through two main dimensions: 1) training the models to generate thinking tokens via large-scale reinforcement learning, and 2) the development of search techniques to assess different reasoning paths at test time. Both approaches increase the amount of inference compute at runtime and latency but yield significant boosts in accuracy and applicability.

The first step requires teaching the model to output an intermediate set of thinking tokens before providing a final answer. These are a sequence of tokens that represent an explicit chain of thought. By training or prompting models to output their reasoning processes, we can guide them in breaking down complex problems into a sequence of intermediate sub-steps. For example, this has been shown to allow a model to correctly answer questions it previously got wrong with direct prompting.⁶ The number of tokens can be controlled at inference to support more nuanced thinking. A classic example is chain-of-thought prompting (CoT). These techniques encourage a model to generate intermediate reasoning steps rather than jump to the answer.

However, simply generating a single chain of thought may not be enough if the model's first attempt contains errors. So, the other dimension is to develop voting and search mechanisms to evaluate the merits of multiple reasoning paths. This provides a way to focus on exploring the most promising ones for generating an answer or response across a tree of possibilities. For example, in tree-of-thought (ToT) frameworks, the model treats partial thoughts as states and uses a search algorithm to find a sequence that leads to a correct answer.

ToT can dramatically improve success on tasks like puzzles and planning by exploring different reasoning branches and even backtracking or revising steps. For example, Princeton and Google DeepMind researchers found that a tree-search approach solved a significantly higher proportion of problems than a single-chain approach.⁷

[6] "Language Models Perform Reasoning via Chain of Thought." Accessed: Apr. 08, 2025. [Online]. Available: <https://research.google/blog/language-models-perform-reasoning-via-chain-of-thought/>

[7] S. Yao et al., "Tree of Thoughts: Deliberate Problem Solving with Large Language Models," Dec. 03, 2023, arXiv: arXiv:2305.10601. doi: 10.48550/arXiv.2305.10601

Creating a foundation for reasoning

Organizing the generation of useful reasoning chains is an active research area using various approaches. Prompting the model to decompose the problem into steps or drawing on specific facts has shown some promise. Other research has explored ways to train models on examples of good reasoning, such as data sets showing intermediate steps. Models can also be trained to annotate their output with justifications and citations to ensure claims are backed by good evidence.

Ongoing research at TII and other groups is exploring ways to improve the efficiency of reasoning since generating many reasoning paths can be computationally expensive and slow, since there is a tradeoff between computational speed and accuracy. Improving the quality of reasoning paths is important since a model may produce lengthy but nonsensical chains. Process supervision methods for training a secondary reward model can improve efforts to judge intermediate steps to prune implausible paths sooner. It is also important to improve techniques for ensuring consistency of reasoning chains with known facts or logical rules, since long chains may drift into hallucination without periodic grounding.

TII's researchers are actively exploring these fronts to create a strong foundation for reasoning. Although this is a developing field, the consensus is that test-time reasoning augmentation is a necessary complement to pretraining. It allows models to transcend their training limitations by thinking for longer and exploring alternatives when faced with novel challenges, which is a critical capability on the path to AGI.

Training better reasoning models

In addition to clever prompting, thinking token generation techniques, and reasoning path search techniques, achieving robust reasoning requires training the models to reason explicitly. New reasoning models use new approaches to train reasoning abilities before providing a final answer. This reasoning ability is not expressly programmed but rather emerges through training processes like reward-based reinforcement learning and rule-based verifiable output reinforcement learning.

These reasoning models can perform some reasoning before providing a final answer. These reasoning abilities are not coded explicitly. Rather, reinforcement learning allows the model to learn reasoning abilities that take advantage of the underlying foundation model. The excitement with RL approaches stems from their ability to automatically discover new reasoning paths that sometimes surpass human capabilities. For example, the reasoning model can generate thinking tokens and explore the merits of different thinking paths before providing a final answer.

Reward-Based Reinforcement Learning involves training a reward model that evaluates the quality of generated answers and then uses this reward model to optimize the main policy or model. These take advantage of policies that define what a better reasoning process might look like. There are many ways to enhance the creation and use of these policies. For example, various policy gradient-based RL methods, such as Proximal Policy Optimization (PPO) and variants like Joint Reward Policy Update (JRPU), require an AI developer to define and train a reward model. This guides the reasoning model in learning better reasoning processes based on the desired behavior. However, this also requires different training techniques than standard foundation model training. TII has published work in this area, including the use of generative flow networks (GflowNets), which aim to generate a more diverse set of solutions compared to standard RL by training the model to sample answers proportional to the reward function.⁸ This approach is still in the early development phase, but could eventually allow reasoning models to discover reasoning techniques that AI developers did not previously consider.

[8] R. Alami, A. Abubaker, M. Achab, M. E. A. Seddik, and S. Lahlou, "Investigating Regularization of Self-Play Language Models," Apr. 04, 2024, arXiv: arXiv:2404.04291. doi: 10.48550/arXiv.2404.04291.

Training better reasoning models

Rule-based RL directly rewards the model based on whether its final answer to a given problem is correct, without relying on a separate reward model or predefined reasoning steps. In this case, the reward is derived from verifiable rules or vetted data sets such as math solutions. These provide an automated objective measure of correctness. This method offers more flexibility for the model to discover new reasoning pathways since it does not rely on complying with existing human approaches. TII relies on this approach since it has the potential to surpass human limitations.

Despite progress, challenges remain in training reasoning models. Balancing multiple objectives (correctness, conciseness, harmlessness, etc.) in RL is difficult – models can regress on one metric while optimizing another. There's also the risk of over-optimization: a model might learn to game the reward (e.g., producing answers that superficially satisfy the reward model but are wrong on careful analysis, a form of reward hacking).

Moreover, RL training is resource-intensive; not every organization has the resources to fine-tune large models for better reasoning. This is why TII's open models and the broader open-source ecosystem are so vital. They allow many researchers to experiment and share the load of improving these systems. By pooling our collective efforts, we can iterate faster on training methods that yield emergent reasoning abilities in AI.

Creating a foundation for ground truth

No matter how good a model's reasoning mechanics are, AGI will remain elusive if it cannot ground its reasoning in true and accurate information. Today's LLMs are prone to hallucination. They produce statements that sound plausible but are factually false precisely because they lack a direct connection to verifiable ground truth data. Building a foundation for ground truth means integrating sources of verifiable data into the model's operation and training such that its outputs can be checked and trusted.

Verifiable data refers to information whose correctness can be confirmed by an external reference or reliable procedure. For example, a historical fact can be verified by consulting an encyclopedia; a mathematical result can be verified by calculation or proof; a logical assertion can be verified by a formal checker or by consistency with known axioms.

In the context of reasoning models, verifiable data often comes in the form of retrieved documents, databases, or API results that the model can reference as it formulates an answer. Unlike the model's internal parameters (which contain a blurred statistical memory of training data), an external source like Wikipedia, a curated knowledge graph, or an API call to a system of record provides a concrete, citable basis for factual claims. Incorporating such data helps ensure that each step of the model's reasoning is anchored to something that can be independently confirmed.

Examples of types of verifiable data include the following:

- **Mathematical and logical solvers:** These tools provide ground truth for well-defined problems. A prime example is a calculator or math engine. If the model needs to compute 26×24 , it can rely on the calculator (which is always correct) rather than its unreliable arithmetic instincts.

Verifiable data is essential for training reasoning models using reinforcement learning. For example, solving math problems consists of a set of prompts of problems and known final solutions. This allows the model to be rewarded for correct answers, even if the training data does not explicitly provide the intermediate reasoning steps, which the model should discover in an unsupervised manner.

TII is actively involved in curating reasoning datasets, including sourcing math tests with problems and solutions. The strategy involves training the model on the distribution of problems, potentially breaking down complex problems into solvable sub-problems to enable generalization and the ability to tackle novel challenges.

To build AGI, we must build AI that is grounded in reality. Verifiable data and grounding techniques are the foundation to ensure our reasoning models evolve into trustworthy problem solvers that know and show the basis for their conclusions.

Extending model capabilities

Achieving AGI will also require extending our models beyond pure language reasoning into the ability to act upon the world and use external tools. A truly general intelligence should not only think but also be able to leverage instruments (whether software tools or physical actuators) to carry out tasks, retrieve information, and interface with the environment. Therefore, an important part of building a better foundation is enabling models to use external tools and APIs, thereby enhancing their capabilities and allowing more agentic behavior.

External tool use refers to the ability of a model to interact with external tools or functions to enhance its capabilities. The foundation model needs to learn how to parse and format it in a way that can be understood and utilized by a specific tool, such as a calculator, search engine, other specialized models, or an application's API.

More broadly, this supports the development of agentic AI systems that can take advantage of various resources to solve tasks. The TII is engaged in pre-processing data to facilitate this, creating datasets of prompts and corresponding tool inputs and outputs.

Modern language models are increasingly being designed as platforms for tool use. In practical terms, this means a model can decide to call a calculator API for arithmetic, query a search engine for information, invoke a translation service for converting languages, or execute code to solve a programming task.

Types of tools that an AI agent might use include, but are not limited to:

- **Calculators and Math Solvers:** for arithmetic, algebra, calculus – any exact computation.
- **Search Engines and Databases:** for retrieving up-to-date information or specific facts from the web or a knowledge base.
- **APIs for services,** e.g., weather APIs, flight booking APIs, maps and navigation APIs, etc., to fetch structured information or perform transactions.
- **Code Interpreters:** the model writes and executes code (Python, for instance) to solve a problem or simulate something. This has been used in models like GPT-4 for complex math or data analysis tasks.

Integrating these tools turns a static model into an AI agent that can interact with the world. For example, an agent tasked with answering a complex question might query a search engine for relevant articles, then ask a calculator to crunch some numbers from those articles, and then compose a final answer. The model orchestrates this chain of actions through a reasoning policy.

Extending model capabilities

While tool use greatly extends capability, it introduces challenges that need further research:

- **Tool selection and planning:** The model has to decide which tool to use and when. Using tools effectively is a sequential decision problem. The agent might have many possible actions at each step (e.g., multiple API calls it could make). Planning a sequence of tool uses (like a mini-program) to solve a task is non-trivial. This is an area where combining search/planning algorithms with learned policies is being explored.
- **Understanding tool outputs:** The model must be able to interpret the result returned by a tool and incorporate it into its context. This might involve reading a JSON response from an API or the text output of a calculator and then adjusting its subsequent reasoning. Ensuring the model accurately parses and uses these outputs is essential.
- **Error handling and robustness:** Tools can fail or return unexpected results. An advanced agent should detect when an API returns an error or irrelevant info and try an alternative approach. This kind of robustness requires the model to have some meta-cognitive ability to assess the success of an action.
- **Security and safety:** An agent that can execute code or make API calls could also yield to harmful behavior if not constrained (e.g., call a destructive API, leak private info via a tool, etc.). Research into sandboxing AI actions and preventing misuse is crucial. Tools should ideally have permissioning, and the AI should be trained to follow usage policies.
- **Resource use and efficiency:** Different types of tools come with various time and cost tradeoffs. If a model overuses tools for trivial things, it could slow down responses. There's a need for the agent to learn when it truly needs a tool and the costs and benefits of various tools to balance efficiency.

TII is actively working on enabling tool use in its models. A notable effort is preprocessing and dataset creation for tool augmentation that shows models how to invoke tools. By including many such examples, models such as the Falcon 3 series can learn to follow that pattern to learn the appropriate API calling syntax for different tools.

Extending model capabilities through tool use is an essential step toward AGI because it transforms passive text predictors into active problem-solvers that can interface with the real world. As we give models the ability to act in a way constrained by reasoning and grounded by verifiable data, we inch closer to AI that can autonomously perform complex tasks.

The synergy of a strong foundation model, explicit reasoning, grounding in truth, and tool-augmented action is what will eventually yield agentic AI systems far more capable than the sum of their parts. Each aspect addresses a fundamental limitation of current models, and together, they form the outline of an AGI architecture in development.

Conclusion

TII's research in foundational models, reasoning, reinforcement learning, verifiable data, and external tool use represents significant contributions to the broader pursuit of AGI. Its work on open-source models can foster collaboration and allow the wider research community to build upon these advancements.

Continued research and collaboration in these fundamental areas are crucial to realizing more advanced and capable artificial intelligence systems. Future efforts will focus on scaling foundation models, exploring the merits of novel model architectures, enhancing reasoning capabilities, and exploring innovative training methodologies like generative flow networks to push the boundaries of AI.

The journey to AGI requires a multi-pronged approach, and the Technology Innovation Institute is contributing across all these fronts to build a better foundation. Scaling LLMs alone is not enough. Instead, we must combine strong foundation models with advanced reasoning techniques, grounding in truth, and tool-use capabilities. By contributing these resources to the open community, TII follows the belief that AGI should be a shared endeavor built transparently and collaboratively.

Building a better foundation for AGI is an ambitious goal that is being tackled step by step. We need excellence in our foundation models, plus the glue that allows them to reason, verify, and act. The progress so far, from LLMs that can solve math problems with explanations to agents that can browse the web, gives a glimpse of an emerging new breed of AI.

TII's work, alongside global research efforts, is knitting together these advancements into a cohesive path forward. We encourage developers and researchers everywhere to take advantage of the open models and tools discussed, to experiment with new ideas (for example, try fine-tuning a Falcon model with your own chain-of-thought data or integrating it with a new API), and to share feedback and results.

Every insight helps. By iterating together on better reasoning, grounding, and extensibility, we can collectively build on the foundation laid so far and accelerate progress toward safe, reliable, and truly general AI.