

Navigate AI from mission to measurable impact

A practical framework for responsible AI adoption in child welfare

Child welfare agencies are under pressure to determine where AI can meaningfully improve practice, and where it should not be used. The challenge is not simply choosing technology. It is identifying problems worth solving, protecting professional judgment, understanding risk, engaging the workforce, and generating evidence that a solution works before scaling it. Generative AI could help child welfare agencies reduce administrative burden, make trusted information easier to use, and give staff more time to focus on children and families. But agencies need to identify which child welfare decisions should remain entirely human, which administrative tasks can be augmented, and what evidence agencies need before expanding AI use to strengthen child welfare practice.

Move from AI possibilities to practical, evidence-informed decisions

Mathematica helps state child welfare agencies move from AI possibilities to defensible adoption decisions, from setting outcomes and identifying promising opportunities through testing, evaluation, and continuous improvement.

Our approach is technology- and vendor-agnostic. We don't start with an AI product. We start with your mission, desired outcomes, and workforce pain points, then determine where AI can add value, where it should not be used, and what must be true before a use case moves forward.

Agencies can use our services individually or combine them into a phased engagement, making it possible to begin with readiness and use-case exploration before committing to a particular technology or large-scale implementation.

Built for the realities of child welfare

For child welfare agencies, AI cannot be separated from practice. Mathematica combines deep child welfare expertise with human-centered design, data and technology expertise, responsible AI governance, implementation support, and independent evaluation. We understand how policy, workflow, professional judgment, workforce conditions, and family experience interact. The goal is not more AI. It is better decisions about where AI can reduce avoidable burden, improve access to trusted information, free capacity for engagement and critical thinking, advance your mission, and meet the safeguards and evidence required to scale.

Depending on your agency's needs, our five-stage framework provides a practical path from exploration to responsible scale:

1. Start with mission and needs.

Work with agency leaders to define priority outcomes, then engage the people closest to the work to identify pain points and promising AI use cases.

2. Navigate the AI landscape.

Determine whether AI is appropriate and, if so, assess options and help decide whether to buy, adapt, configure, or not proceed.

3. Set boundaries and governance.

Establish where AI can assist staff, where human judgment must remain central, and what safeguards must be in place before moving forward.

4. Co-design and test before scaling.

Design and test potential solutions with the workforce, while engaging community members to help ensure proposed use cases are acceptable and build social license for adoption.

5. Measure, learn, and improve.

Support implementation and assess whether AI is working as intended and improving the outcomes that matter to your agency before expanding its use.

Explore what responsible AI could mean for your agency:

Allon Kalisher, Senior Child Welfare Lead, akalisher@mathematica-mpr.com

Elizabeth Weigensberg, Principal Researcher, eweigensberg@mathematica-mpr.com

For more information, visit [Mathematica.org/focus-areas/human-services/child-welfare](https://mathematica.org/focus-areas/human-services/child-welfare)