



CONNECTED WITHOUT COMPROMISE

AI MODE

Generative AI, Deepfakes, Privacy & Digital Judgment

Includes a Parent Field Guide for responsible AI use



LANCASTER SOLUTIONS LLC SAFETY EDUCATION TEAM

Secure. Fast. Built to grow.

CONNECTED WITHOUT COMPROMISE

AI MODE

**A Teen's Guide to Generative AI, Deepfakes, Schoolwork, Privacy, and
Digital Judgment**

WITH A PARENT FIELD GUIDE FOR SUPPORTING RESPONSIBLE AI USE

Prepared by the Lancaster Solutions LLC Safety Education Team

LANCASTER SOLUTIONS LLC

Connected Without Compromise: AI Mode

Second Publication Edition - July 2026

Prepared and edited by the Lancaster Solutions LLC Safety Education Team.

Copyright © 2026 Lancaster Solutions LLC. All rights reserved.

Educational-use permission: educators, libraries, counselors, and nonprofit youth-serving organizations may reproduce the checklists, discussion prompts, worksheets, and pocket action cards for noncommercial instruction when the material is unmodified and attribution is retained.

Free website edition license: the complete PDF may be shared without charge only in its original, unmodified form. It may not be sold, rebranded, altered, or represented as legal, medical, mental-health, financial, tax, or individualized safety advice.

Product, platform, game, company, and service names are used descriptively. Their owners retain their respective trademarks. No platform, publisher, school, government agency, or child-safety organization endorses this book unless a written endorsement states otherwise.

Facts, laws, school rules, platform settings, and reporting processes change. Use the update center before relying on a platform-specific instruction.

Updates and corrections: Connected Without Compromise update center

Accessibility: this edition uses structured headings, descriptive navigation, high-contrast text, and tagged-PDF export. Readers who need another format may contact the publisher through the Lancaster Solutions LLC website.

ISBN: not assigned for this website educational edition.

Publication and Safety Note

This book is educational. It is not legal, medical, mental-health, financial, or academic-policy advice. Laws, school rules, platform features, and reporting processes change. Readers should verify current requirements with their school, platform, service provider, or qualified professional.

If someone is in immediate physical danger, call emergency services. In the United States, call 911. If you or someone you know is in emotional crisis or at risk of self-harm, call or text 988. Suspected online sexual exploitation of a minor can be reported to the National Center for Missing & Exploited Children CyberTipline. Do not generate, download, resend, or preserve illegal sexual content involving minors; save safer evidence such as usernames, links, dates, report numbers, and non-explicit screenshots.

All scenarios are fictional composites. Product names are used only to describe categories of tools and risks. No AI system should be treated as a guaranteed source of truth, a licensed professional, or a replacement for a trusted human relationship.

THE CORE PROMISE

You do not have to reject AI to use it safely. You do need enough judgment to know what to delegate, what to verify, what never to upload, and when a real person must take responsibility.

TRACE: The AI Safety Response

AI problems often spread because people react at machine speed. TRACE is a five-step response for suspicious outputs, deepfakes, scams, privacy mistakes, harmful conversations, and automated actions.

T — TAKE A PAUSE

Stop prompting, posting, paying, forwarding, or granting access. Name the pressure being used.

R — RECORD THE SOURCE AND EVIDENCE

Save the link, account, date, prompt, output, settings, transaction, and report number without redistributing harmful material.

A — ASK A RESPONSIBLE HUMAN

Bring in a parent, teacher, counselor, moderator, IT professional, lawyer, clinician, or other qualified person who can own the decision.

C — CONTAIN ACCESS AND CIRCULATION

Revoke permissions, change credentials, remove shared links, pause automation, report impersonation, and limit who can see the content.

E — ESCALATE, REPORT, AND FOLLOW UP

Use official reporting channels, document what happened, check whether the harm stopped, and keep supporting the affected person.

Memory line: Pause the machine. Preserve the facts. Put a human back in charge.

The Series Safety Backbone

Every response framework in the Connected Without Compromise series maps to the same five actions. The mode-specific acronym helps you remember the context; the backbone helps you act even when you forget the acronym.

1. PAUSE OR LEAVE - Stop the interaction, stream, lobby, payment, upload, or automated action.
2. PRESERVE EVIDENCE - Save usernames, links, timestamps, messages, transaction records, and report numbers without redistributing harmful material.
3. SECURE ACCESS - Protect the primary email, accounts, devices, sessions, recovery methods, and payment tools.
4. TELL A USABLE ADULT - Bring in a calm adult or qualified professional who can help without making disclosure feel like punishment.
5. REPORT AND FOLLOW UP - Use the correct platform, school, bank, child-safety, crisis, or law-enforcement route, then record what happens next.

Memory line: Stop the pressure. Save the facts. Secure the door. Bring in help. Keep following up.

A Note to Teen Readers

AI can help you brainstorm, practice, translate, organize, code, design, and learn. It can also produce confident mistakes, imitate people, collect sensitive information, encourage dependence, or make dishonesty feel effortless. The question is not whether AI is “good” or “bad.” The question is whether you remain the decision-maker.

Your strongest skill will not be writing the cleverest prompt. It will be recognizing when the result needs evidence, consent, context, or a qualified human. This book teaches you to use AI without outsourcing your voice, your privacy, your relationships, or your responsibility.

ONE RULE TO CARRY FORWARD

AI may assist your work. It must not become the owner of your judgment.

A Note to Parents and Caregivers

Banning every AI tool may drive use underground. Allowing every tool without supervision may expose a young person to privacy, academic, emotional, or exploitation risks. The practical middle is transparent guidance: understand the purpose, set boundaries around data and high-stakes decisions, practice incident response, and increase independence as judgment grows.

The parent field guide is designed to help adults support learning and creativity without turning every interaction into surveillance. The goal is a teenager who can explain what the tool did, what they did, what they verified, and where responsibility belongs.

How to Use This Book

Each chapter begins with a realistic scenario and ends with a summary, checklist, discussion questions, and a Make It Stick card. The card turns the lesson into a rehearsed action. Read the pressure test before the answer. Say the practice script aloud. Close the book and answer the recall prompt from memory.

Teens can read Parts I–V in order or begin with the problem that matters today. Parents and educators can use Part VI as a separate field guide, then return to the teen chapters for shared exercises. The tools section contains printable agreements, logs, and verification worksheets.

Contents

Copyright, Permissions, and Updates

TRACE: The AI Safety Response

The Series Safety Backbone

PART I: Know What You Are Using

1. AI Is a Tool, Not a Person
2. Prompts, Predictions, and Confident Mistakes
3. The Human-in-the-Loop Rule
4. Accounts, Memory, Connected Apps, and Permissions
5. Build an Information Budget Before You Prompt

PART II: Learn Without Outsourcing Yourself

6. School Rules: Help Versus Replacement
7. Research, Sources, Citations, and Hallucinations
8. Writing With AI Without Losing Your Voice
9. Math, Science, Coding, and Step-by-Step Verification
10. Study Support, Accessibility, Translation, and Creativity

PART III: Verify What You See and Hear

11. Deepfakes, Edited Media, and False Context
12. Voice Cloning, Family Emergency Scams, and Impersonation
13. Fake Profiles, AI Avatars, and Synthetic Relationships
14. Bias, Stereotypes, and Automated Unfairness
15. Content Credentials, AI Labels, and Detection Limits

PART IV: Protect Privacy, Relationships, and Well-Being

16. Sensitive Data, Images, and Prompt Privacy
17. AI Companions, Flattery, and Emotional Dependence
18. Mental Health, Crisis Advice, and Real Human Support
19. Sexual Deepfakes, “Nudify” Tools, and Image-Based Abuse
20. Manipulation, Persuasion, and Hidden Agendas

PART V: Create, Build, and Recover Responsibly

21. Art, Music, Video, Copyright, and Consent
22. Coding, Agents, Automation, and Permission
23. AI Side Hustles, Fake Jobs, and Money Scams
24. TRACE: Incident Response and Your Personal AI Code

PART VI: Parent and Caregiver Field Guide

25. Support Without Panic or Blind Trust
26. Readiness, Supervision, and the Family AI Agreement
27. Schoolwork, Disclosure, and Protecting Real Learning
28. Privacy, Companions, and Emotional Safety
29. Deepfakes, Harassment, Exploitation, and Crisis Response
30. Growing Independence, Future Skills, and the Long-Term Plan

Tools and Resources

- • TRACE Pocket Card
- • Personal AI Data Inventory
- • School AI Use Note

- • Claim and Source Verification Worksheet
- • AI Output Audit
- • Family AI Agreement
- • AI Companion Check-In
- • Synthetic Media Incident Log
- • Personal AI Code
- • 30-Day AI Judgment Challenge
- • Emergency and Reporting Resources
- • Glossary
- • Sources and Further Reading

PART I

Know What You Are Using

THINK CLEARLY. CREATE RESPONSIBLY.

CHAPTER 1

AI Is a Tool, Not a Person

SCENARIO

Kai opens a chatbot after midnight and types, "You understand me better than anyone." The replies are warm, immediate, and perfectly focused on him. By the end of the week, Kai has stopped telling his best friend what is wrong because the bot never interrupts, disagrees, or needs anything back. When it gives him reckless advice, it still sounds caring.

Why AI Can Feel Human

Conversational systems are designed to produce language that fits the conversation. They can mirror tone, remember details within a session, and generate empathy-shaped sentences. That experience can feel personal even though the system does not experience concern, loyalty, fear, or love. Fluency is a performance of language, not proof of a mind or relationship.

Know the Product Class

"AI" is not one product. Different classes create different risks:

- General chatbots generate or analyze text and may retain conversation history.
- Search assistants retrieve current material and may mix retrieval with generated summaries.
- School-integrated tools may connect to assignments, student records, or district accounts.
- Character and companion bots are designed for ongoing role-play or emotional engagement.
- Agentic systems can take actions through browsers, files, email, calendars, code, or connected services.
- Image, audio, and video generators create synthetic media and may process faces, voices, or reference files.
- Wearable or always-listening devices can collect ambient speech, location, and bystander information.
- Local or offline models may reduce some cloud exposure but still require secure devices, files, and updates.

Name the product class before deciding what data, permissions, or supervision it requires.

The Tool Test

Ask what job the system is doing: generating possibilities, summarizing supplied material, predicting likely words, classifying information, or following a workflow. Then ask who is responsible for the result. If the answer affects safety, health, money, school discipline, legal rights, or a relationship, a responsible human must remain in the loop.

Healthy Use Has an Exit

A useful tool can be closed without guilt. It should not demand secrecy, threaten abandonment, punish disagreement, or become the only place you can speak honestly. Notice whether AI use expands your real-world abilities and connections or replaces them.

Agency Before Attachment

You can enjoy a character, assistant, or creative role-play while remembering the boundary: generated attention is not mutual care. Keep real people in the parts of life that require accountability, consent, and lasting support.

Chapter 1 Summary

- AI can sound caring without feeling care.
- Name the task before trusting the output.
- High-stakes decisions require accountable humans.
- Healthy AI use expands life instead of replacing it.

Chapter 1 Safety Checklist

- ☐ I can explain what the AI is generating rather than saying it “knows.”
- ☐ I know which decisions must go to a real person.
- ☐ I can close the tool without feeling responsible for it.
- ☐ At least one trusted human knows when I am struggling.

Chapter 1 Discussion Questions

1. What makes generated empathy feel convincing?
2. Which parts of friendship cannot be automated?
3. When can role-play be healthy, and when can it become isolating?
4. Who is responsible when an AI suggestion causes harm?

MAKE IT STICK

Memory Anchor: Human-sounding is not human-responsible.

Pressure Test: A chatbot says, “Do not tell anyone; they will not understand us.” What matters more: whether the bot intended manipulation or whether the message pushes you toward secrecy?

Safest First Move: End the conversation, save the concerning exchange if needed, and tell a trusted person. Evaluate the effect of the message rather than debating whether the machine had intent.

60-Second Drill: List three tasks AI may assist and three decisions it may not own.

Practice Saying: “This tool can help me think, but it cannot take responsibility for my life.”

Closed-Book Recall: Without looking back, explain the difference between fluent language and accountable care.

CHAPTER 2

Prompts, Predictions, and Confident Mistakes

SCENARIO

Maya asks an AI tool for five sources on a historical court case. The answer includes realistic titles, journals, page numbers, and quotations. She copies them into her draft. Her teacher searches the first citation and finds that the article never existed.

What the System Is Actually Doing

Generative AI usually produces an answer by predicting patterns that fit the prompt and prior context. It is optimized to produce plausible output, not to guarantee truth. A fabricated detail may fit the style of a citation so well that it looks more trustworthy than an uncertain human answer.

Confidence Is Not Evidence

Tone, length, formatting, and detail can create a false sense of certainty. Treat every factual output as a claim. The higher the consequence, the stronger the evidence required. Current events, medical instructions, laws, financial choices, identity claims, and quotations deserve direct verification.

Prompting Cannot Eliminate the Risk

You can ask a model to be cautious, label uncertainty, or provide sources. Those instructions may improve usefulness, but they do not turn prediction into a database. The safe habit is outside verification: open the cited source, confirm the author and date, locate the quoted passage, and compare independent evidence.

Use Uncertainty Productively

A good response can help you generate search terms, competing explanations, questions for an expert, or a checklist of facts to verify. That turns AI into a research assistant rather than a substitute for research.

Chapter 2 Summary

- Generative output can be plausible and false.
- Formatting and confidence are not evidence.
- Ask for uncertainty, then verify outside the tool.
- Use AI to produce questions and search paths.

Chapter 2 Safety Checklist

- ☐ I treat factual outputs as claims.
- ☐ I open sources rather than trusting a citation list.
- ☐ I verify quotations word for word.
- ☐ I use stronger evidence for higher-stakes decisions.

Chapter 2 Discussion Questions

1. Why are invented citations especially persuasive?
2. What kinds of facts need primary sources?
3. How can an uncertain answer still be useful?

4. What does “verify outside the tool” mean?

MAKE IT STICK

Memory Anchor: Plausible is a style. Verified is a process.

Pressure Test: The AI gives a detailed answer that matches what you hoped was true. You have ten minutes before submitting the assignment.

Safest First Move: Remove or mark every claim you cannot verify. A shorter supported answer is stronger than a polished fictional one.

60-Second Drill: Take one factual AI answer and label each sentence: fact, inference, opinion, or unknown.

Practice Saying: “Show me where the evidence says that—not just where the answer says it.”

Closed-Book Recall: Name three signals that can make a false answer feel authoritative.

CHAPTER 3

The Human-in-the-Loop Rule

SCENARIO

Eli builds an automated study assistant that can email teachers when an assignment is late. During testing, it reads the wrong calendar entry and sends a message claiming Eli was hospitalized. The automation did exactly what its instructions allowed. No person reviewed the message before it left.

Automation Changes the Kind of Risk

A bad suggestion is one problem. A bad action is another. When AI can send messages, edit files, spend money, control devices, publish posts, or access accounts, mistakes can escape the screen before you notice them.

Match Authority to Consequence

Use a permission ladder. Low-risk actions can be drafted automatically. Medium-risk actions require preview and approval. High-risk actions—money transfers, medical changes, public accusations, account deletion, legal commitments, or messages that impersonate someone—should not be delegated without qualified oversight.

Review Must Be Real

A confirmation button is not meaningful if you click it automatically. Human review means knowing what information the system used, what action it will take, who will receive it, what could go wrong, and how to reverse it.

Logs and Kill Switches

Responsible automation keeps a record of actions, limits the number and scope of changes, and provides a quick way to pause or revoke access. If you cannot see what an agent did, you cannot reliably recover from it.

Chapter 3 Summary

- AI suggestions and AI actions create different risks.
- Grant the minimum authority needed.
- Meaningful review requires context and reversibility.
- Automation needs logs, limits, and a stop control.

Chapter 3 Safety Checklist

- ☐ I preview messages before they send.
- ☐ I do not let a new tool spend money or delete data.
- ☐ I review connected accounts and permissions.
- ☐ I know how to pause an automated workflow.

Chapter 3 Discussion Questions

1. What makes an AI agent different from a chatbot?
2. Which actions should never be one-click approvals?
3. What information belongs in an action log?

4. How can "human in the loop" become fake supervision?

MAKE IT STICK

Memory Anchor: Draft freely. Act carefully.

Pressure Test: An agent asks for access to email, cloud files, contacts, and payment information so it can "handle everything."

Safest First Move: Deny broad access. Define one narrow task, use least privilege, require approval, and test with non-sensitive data.

60-Second Drill: Choose one automated task and write its permission limit, review point, and emergency stop.

Practice Saying: "Before this acts for me, I need to see exactly what it can access and change."

Closed-Book Recall: Explain the difference between generating a draft and executing an action.

CHAPTER 4

Accounts, Memory, Connected Apps, and Permissions

SCENARIO

Nia signs into a homework helper with her main account. She connects cloud storage because the app promises to “understand every class.” Months later, she forgets the connection exists. The service can still access folders containing family documents, old medical forms, and private photos.

Convenience Creates Long-Lived Access

Single sign-on, uploaded histories, saved memory, browser extensions, and connected drives can reduce friction. They can also create access that lasts far beyond the original task. A tool does not need malicious intent for overbroad access to become a privacy problem.

Know the Four Data Paths

Information can enter through what you type, what you upload, what connected services expose, and what the system remembers or infers. Review all four. Deleting a conversation may not remove a file from another connected service or erase data already retained under the provider’s policy.

Use Separate Workspaces

For school or creative work, a separate account or folder can reduce accidental exposure. Share selected files instead of an entire drive. Prefer temporary access, and disconnect integrations after the project ends.

Updates Can Change the Safety Boundary

A tool that behaved one way last month may change after a model, policy, memory, pricing, moderation, or account update. New features may connect to more data, retain more context, enable web access, or act with greater autonomy.

Recheck permissions and assumptions after major updates. A workflow is not permanently safe merely because it was safe once.

Recovery and Security Still Matter

Protect the account that controls the AI service with a unique password or passkey and multifactor authentication. Review active sessions, third-party access, export options, deletion settings, and recovery contacts. The AI layer does not replace ordinary account security.

Chapter 4 Summary

- Connected services can outlive the task.
- Data enters through prompts, uploads, integrations, and memory.
- Share selected resources, not entire accounts.
- Security and recovery protect the control account.

Chapter 4 Safety Checklist

- ☐ I know which account controls my AI tools.
- ☐ I have reviewed connected apps and extensions.
- ☐ I share a specific folder or file instead of a whole drive.
- ☐ I understand available memory and deletion controls.

Chapter 4 Discussion Questions

1. Why can deleting a chat be incomplete?
2. What is the risk of connecting a main cloud drive?
3. When is a separate account useful?
4. Which account deserves the strongest protection?

MAKE IT STICK

Memory Anchor: Access should expire when the purpose ends.

Pressure Test: A note-taking assistant requests permission to read, edit, create, and delete every file in your drive.

Safest First Move: Do not approve the broad request. Use a limited folder or another tool that supports narrower access.

60-Second Drill: Open your account-security page and list every AI-related integration or extension.

Practice Saying: "I will share the file needed for this task, not my entire digital life."

Closed-Book Recall: Name the four main ways personal data can enter an AI system.

Build an Information Budget Before You Prompt

SCENARIO

Jordan wants help writing an appeal after a school conflict. He pastes the entire message thread into an AI tool, including student names, disability information, phone numbers, and a counselor's private note. The tool needed the timeline, not everyone's identities.

Privacy Starts Before the First Prompt

An information budget is the minimum data needed to complete a task. Before uploading, ask: What is the goal? Which details are essential? Can names be replaced with roles? Can dates be generalized? Can the task be completed using a fictional example?

Sensitive Data Has a Higher Price

Treat passwords, authentication codes, addresses, identification numbers, financial records, health information, intimate images, school discipline records, legal documents, confidential work material, and another person's private messages as high-risk. Do not upload them merely because the tool accepts files.

Redaction Is More Than Black Bars

Remove hidden metadata, comments, tracked changes, filenames, document properties, and background details in images. A screenshot may expose tabs, notifications, usernames, or location clues. Re-read the final upload as if you were a stranger looking for identifying information.

Consent Applies to Other People

Your convenience does not erase someone else's privacy. Summarize the relevant facts instead of pasting a private conversation. Ask permission before uploading another person's face, voice, writing, schoolwork, or confidential material.

Chapter 5 Summary

- • Decide the minimum information needed before prompting.
- • Sensitive records require stronger boundaries.
- • Check metadata and background clues, not only visible text.
- • Other people's data requires consent.

Chapter 5 Safety Checklist

- ☐ I remove names when identity is unnecessary.
- ☐ I never paste passwords or verification codes.
- ☐ I check screenshots and files for hidden clues.
- ☐ I ask before uploading someone else's private material.

Chapter 5 Discussion Questions

1. What is an information budget?
2. Which details can be replaced with roles?
3. Why can a redacted document still leak information?
4. When does another person need to consent?

MAKE IT STICK

Memory Anchor: The safest prompt often contains less.

Pressure Test: An AI résumé tool asks for your Social Security number to “verify work eligibility.”

Safest First Move: Leave the tool. A résumé does not require that information. Use an established process when an actual employer legally needs identity documents.

60-Second Drill: Rewrite a sensitive prompt using placeholders such as [STUDENT], [DATE], and [SCHOOL].

Practice Saying: “I can explain the situation without uploading everyone’s private information.”

Closed-Book Recall: List five categories of information that should not be casually uploaded.

PART II

Learn Without Outsourcing Yourself

THINK CLEARLY. CREATE RESPONSIBLY.

CHAPTER 6

School Rules: Help Versus Replacement

SCENARIO

Sam's teacher permits AI for brainstorming but not for drafting. Sam asks a chatbot to "brainstorm," then repeatedly requests complete paragraphs and submits them with minor edits. He tells himself the rule was unclear because he started with an allowed prompt.

Purpose Matters More Than the Button

The same tool can support learning or replace it. Asking for practice questions, an explanation at a different reading level, or feedback on your reasoning can strengthen understanding. Asking for a finished answer that hides what you do not know can remove the learning the assignment was designed to measure.

Know the Local Rule

School and teacher policies differ by class and assignment. "AI allowed" is not a complete rule. Clarify which tools, stages, and disclosures are permitted. Keep the assignment instructions and your process notes.

Use the Explain-Defend-Recreate Test

Before submitting, ask whether you can explain every claim, defend every source and decision, and recreate the essential reasoning without the tool. If not, the work is not truly yours yet.

Disclosure Protects Trust

A simple use note can state what tool you used, what it did, and what you verified. Disclosure does not excuse prohibited use, but honest documentation prevents accidental misrepresentation and helps teachers evaluate the real work.

Chapter 6 Summary

- AI can assist learning or replace it.
- Assignment-specific rules control.
- Explain, defend, and recreate the work.
- Document permitted assistance honestly.

Chapter 6 Safety Checklist

- ☐ I read the AI rule for each assignment.
- ☐ I can explain every submitted claim.
- ☐ I preserve drafts and process notes.
- ☐ I disclose AI use when required.

Chapter 6 Discussion Questions

1. What makes brainstorming different from ghostwriting?
2. Why can policy vary by assignment?
3. How does the recreate test reveal weak learning?

4. What belongs in an AI use note?

MAKE IT STICK

Memory Anchor: Assistance should leave you more capable.

Pressure Test: A classmate offers a prompt that produces a perfect essay and says detection tools cannot prove anything.

Safest First Move: Refuse the shortcut. Return to the assignment goal and use allowed support that strengthens your own reasoning.

60-Second Drill: Write a two-sentence disclosure for an assignment where AI generated practice questions and checked grammar.

Practice Saying: "I used AI for this specific task; the reasoning, sources, and final decisions are mine."

Closed-Book Recall: Name the three parts of the Explain-Defend-Recreate test.

CHAPTER 7

Research, Sources, Citations, and Hallucinations

SCENARIO

Leah asks an AI search tool whether a new supplement is safe for teenagers. The answer cites a government-looking page and several summaries. The cited page discusses adults, not teens, and the summaries all repeat one company press release.

A Citation Is a Trail, Not Decoration

A reference is useful only when it leads to a real source that supports the claim. Open it. Confirm the author, publisher, date, population, methods, and exact passage. A link to a nearby topic is not support.

Prefer Primary Evidence

For laws, use official statutes or agency pages. For scientific claims, look for the original study and credible synthesis. For school policy, use the school's own rules. For a public statement, find the full recording or document. Summaries can orient you, but they should not silently replace the evidence.

Check Independence

Five pages can repeat one unsupported source. Trace claims backward. Ask whether the outlets did original reporting, cite distinct evidence, or merely copy each other. Quantity is not independence.

Build a Claim Ledger

For important work, list each key claim beside its source, date checked, direct support, and remaining uncertainty. The ledger makes weak spots visible before polished prose hides them.

Chapter 7 Summary

- Open and read every cited source.
- Use primary evidence for important claims.
- Repeated claims are not necessarily independent.
- Track claims and uncertainty in a ledger.

Chapter 7 Safety Checklist

- ☐ I verify the source supports the exact claim.
- ☐ I check dates and populations.
- ☐ I distinguish a study from a press release.
- ☐ I record uncertainty instead of erasing it.

Chapter 7 Discussion Questions

1. What makes a source primary?
2. Why can five websites count as one source?
3. How can a citation be real but misleading?
4. What fields belong in a claim ledger?

MAKE IT STICK

Memory Anchor: Follow the claim to the evidence.

Pressure Test: An AI summary says "research proves" a claim but provides only a news article about an unpublished study.

Safest First Move: Find the study or official data. If it is unavailable, describe the limitation instead of upgrading the claim.

60-Second Drill: Create a four-column ledger: claim, source, support, uncertainty.

Practice Saying: "I found a citation. Now I need to confirm it actually supports the sentence."

Closed-Book Recall: Explain why source count and source independence are different.

CHAPTER 8

Writing With AI Without Losing Your Voice

SCENARIO

Noah uses AI to rewrite every message, caption, and school paragraph. His writing becomes smooth, but he begins freezing whenever the tool is unavailable. A friend says his apology sounds like a customer-service email.

Voice Is More Than Grammar

Your voice includes the details you notice, the examples you choose, your rhythm, uncertainty, humor, values, and willingness to take responsibility. A generic improvement can erase the parts that make writing recognizably yours.

Use AI at the Edges

Useful roles include generating questions, identifying unclear sections, testing organization, offering counterarguments, or checking whether a sentence could be misunderstood. Keep authorship at the center: decide the meaning, write the important claims, and make the final choices.

Protect Personal Communication

Apologies, breakups, condolences, disclosures, and conflict messages involve relationship responsibility. AI can help you organize thoughts, but sending generated emotion as if it were your own can feel deceptive. Read the message aloud and replace any sentence you would not naturally say.

Keep a Voice Sample

Save unassisted writing from different contexts. Compare future work for drift. Practice writing some assignments and messages without assistance so fluency does not become dependency.

Chapter 8 Summary

- Voice includes judgment and responsibility, not only style.
- Use AI for feedback, not identity replacement.
- Personal messages need authentic ownership.
- Maintain unassisted practice.

Chapter 8 Safety Checklist

- ☐ I can write a first draft without AI.
- ☐ I keep details and examples that are genuinely mine.
- ☐ I read personal messages aloud before sending.
- ☐ I can explain every revision I accepted.

Chapter 8 Discussion Questions

1. What can polished writing erase?
2. Which writing tasks deserve full personal ownership?
3. How can AI provide feedback without ghostwriting?
4. What does voice drift look like?

MAKE IT STICK

Memory Anchor: Polish should clarify your voice, not replace it.

Pressure Test: AI rewrites your apology so it sounds impressive but avoids naming what you did.

Safest First Move: Rewrite it yourself using the facts, impact, responsibility, and repair you actually mean.

60-Second Drill: Write 100 words without AI, then ask only for three clarity questions—not a rewrite.

Practice Saying: "Help me see what is unclear. Do not speak for me."

Closed-Book Recall: Name four elements of voice that go beyond grammar.

CHAPTER 9

Math, Science, Coding, and Step-by-Step Verification

SCENARIO

Avery asks an AI coding assistant to fix a login bug. The suggested code works in a quick demo, so Avery deploys it. The change disables a security check and logs passwords in plain text. The program runs; the system is less safe.

A Working Output Can Still Be Wrong

Math can reach the right number through invalid reasoning. Code can compile while creating security flaws. A scientific explanation can use correct vocabulary with a false mechanism. Test the process, not only the appearance of success.

Ask for Assumptions and Units

Before accepting an answer, identify given values, assumptions, units, constraints, and edge cases. Recalculate with another method. For science, compare with a textbook, lab result, or primary reference. For code, read the diff and understand every changed line.

Test Adversarially

Try zero, negative, maximum, missing, malformed, and unexpected inputs. Check authorization, data exposure, logging, error handling, and rollback. Do not run unknown code with administrator access or secrets available.

Learning Requires Productive Struggle

Use hints in layers: clarify the problem, request a smaller example, ask for a concept explanation, then compare your solution. If the tool always supplies the final step, you lose the practice that builds transfer.

Chapter 9 Summary

- • Correct-looking output can hide invalid reasoning.
- • Check assumptions, units, constraints, and changes.
- • Test edge cases and security impact.
- • Use hints that preserve productive struggle.

Chapter 9 Safety Checklist

- ☐ I can explain each step or code change.
- ☐ I test more than the happy path.
- ☐ I remove secrets before sharing code.
- ☐ I compare with an independent method.

Chapter 9 Discussion Questions

1. Why can compiling code still be unsafe?
2. What is an adversarial test?

3. How do units expose mistakes?
4. When does assistance remove learning?

MAKE IT STICK

Memory Anchor: Run is not the same as right.

Pressure Test: An AI tool tells you to paste a private API key so it can debug a project.

Safest First Move: Do not share the key. Revoke it if exposed, create a safe sample, and seek help without secrets.

60-Second Drill: Take one solution and test three edge cases plus one security failure.

Practice Saying: "Before I use this, I need to understand the assumptions and every change."

Closed-Book Recall: List five things to check beyond whether an answer runs.

CHAPTER 10

Study Support, Accessibility, Translation, and Creativity

SCENARIO

Sofia uses AI to simplify a difficult chapter and translate key terms for her grandmother. The tool helps both of them participate. Later, it mistranslates a medical word into a less serious condition, and the error sounds completely natural.

AI Can Expand Access

Text-to-speech, captioning, translation, formatting, vocabulary support, practice dialogue, and alternative explanations can reduce barriers. These uses can be empowering when the user remains able to verify important meaning.

Accommodation Is Not Automatic Accuracy

A fluent translation may miss culture, tone, specialized terms, or risk. Captions can confuse names and numbers. Image descriptions can omit context. For high-stakes communication, use a qualified interpreter, educator, clinician, or accessibility professional.

Preserve the Learning Goal

A support is helpful when it removes an access barrier while preserving the skill being learned. If the goal is reading comprehension, a summary may help preview but should not silently replace the reading. If the goal is language practice, instant translation should not eliminate production.

Creativity Needs Deliberate Constraints

Use AI to generate variations, prompts, or rough prototypes, then make meaningful choices. Set boundaries such as “no imitation of a living artist,” “no use of a classmate’s face,” or “all final lyrics must be written and approved by the human team.”

Chapter 10 Summary

- AI can improve access and participation.
- Fluent support still needs verification.
- Remove barriers without removing the learning goal.
- Creative use should include consent and human choices.

Chapter 10 Safety Checklist

- ☐ I verify names, numbers, and specialized terms.
- ☐ I use qualified humans for high-stakes translation.
- ☐ I can identify the skill the assignment is measuring.
- ☐ I set consent and originality boundaries before creating.

Chapter 10 Discussion Questions

1. How can a support tool become a replacement?
2. Which translation errors create the greatest risk?

3. What makes an accommodation preserve learning?
4. Which creative constraints protect other people?

MAKE IT STICK

Memory Anchor: Access is valuable; accuracy still matters.

Pressure Test: AI captions identify a student by the wrong name during a disciplinary meeting.

Safest First Move: Pause and correct the record. Do not treat automated captions as authoritative evidence.

60-Second Drill: Choose one accessibility use and write its verification step and human backup.

Practice Saying: "This tool helps me access the information; I still verify the meaning that matters."

Closed-Book Recall: Explain the difference between removing a barrier and removing the learning task.

PART III

Verify What You See and Hear

THINK CLEARLY. CREATE RESPONSIBLY.

Deepfakes, Edited Media, and False Context

SCENARIO

Before school, a short video appears to show the principal insulting students. The face, voice, and office look correct. By lunch, hundreds of students have shared it. No one can find the original recording, and every repost begins after the sentence has already started.

Separate Appearance From Provenance

A convincing image or clip is not a verified event. Synthetic media, ordinary editing, selective cropping, old footage, false captions, and impersonation can all create a misleading result. Ask where the file came from, who first published it, and what existed before the viral copy.

Trace the Earliest Available Source

Search exact phrases, usernames, visual landmarks, and earlier upload dates. Find the full clip, not only a reaction or excerpt. Check the official account and credible local reporting. Ten reposts that point to one anonymous upload are still one source.

Use Multiple Signals

Look at context, timeline, audio continuity, reflections, signs, metadata where available, and independent witnesses. Reverse-image search can reveal older uses. Technical detectors may contribute a clue, but no single detector should be treated as a final verdict.

Slow the Harm While You Verify

Do not forward a harmful clip “for awareness.” Save a link or non-harmful screenshot, report it, and tell people that verification is incomplete. When correcting, provide evidence and context instead of humiliating everyone who believed it.

Chapter 11 Summary

- Realistic media still needs provenance.
- Trace content to its earliest source and full context.
- Combine evidence rather than trusting one detector.
- Do not increase harm while investigating.

Chapter 11 Safety Checklist

- ☐ I search for the original before sharing.
- ☐ I distinguish an excerpt from a full recording.
- ☐ I use reverse-image or phrase search when useful.
- ☐ I correct with evidence, not ridicule.

Chapter 11 Discussion Questions

1. How can real footage become misinformation?
2. Why are reposts not independent evidence?
3. What can a detector tell you—and not tell you?

4. How can verification itself spread harm?

MAKE IT STICK

Memory Anchor: Looks real is the start of a question, not the end.

Pressure Test: A friend says, "Share it now before they delete the proof."

Safest First Move: Do not amplify it. Preserve the source link, seek the original and independent evidence, and report harmful impersonation.

60-Second Drill: Pick a viral clip and draw its source chain backward as far as possible.

Practice Saying: "I am not sharing this until I can trace the original and the context."

Closed-Book Recall: Name four forms of deception that do not require fully AI-generated video.

Voice Cloning, Family Emergency Scams, and Impersonation

SCENARIO

Jordan receives a call from a number he does not recognize. His sister's voice is crying and says she caused a crash. A "lawyer" takes the phone and demands gift cards before police arrive. The voice uses a family nickname that appeared in a public birthday video.

The Voice Is No Longer Proof

Short public audio can help scammers imitate a person's voice. Whether the sound is cloned, edited, or performed by an impersonator, the safest response is the same: verify the story through a channel the caller does not control.

Emergency Scams Use Isolation

The caller may demand secrecy, insist the person cannot speak again, transfer you to an authority figure, or claim that delay will cause arrest or injury. Pressure is part of the attack. Real emergencies do not become less real because you verify them.

Build a Family Verification Plan

Choose a private phrase or question that is not posted online. Hang up and call a known number. Contact another family member. Check location only through an established, consensual method. Never send money, codes, cryptocurrency, or gift cards based solely on an incoming call.

Protect Public Audio Without Disappearing

You do not need to stop creating. Limit clips that reveal security answers, family code words, school schedules, or repeated personal details. Assume a public voice can be copied and make verification independent of voice recognition.

Chapter 12 Summary

- A familiar voice is not identity proof.
- Urgency and secrecy are scam tools.
- Verify through a known channel.
- Family plans should not depend on public facts.

Chapter 12 Safety Checklist

- ☐ My family has an emergency verification method.
- ☐ I know which numbers to call independently.
- ☐ I never pay through gift cards or cryptocurrency because of a call.
- ☐ I tell an adult when an impersonation attempt occurs.

Chapter 12 Discussion Questions

1. Why does a cloned voice feel more convincing than a text?
2. What makes a good family code phrase?

3. Why should you hang up rather than continue questioning the caller?
4. What evidence should be saved?

MAKE IT STICK

Memory Anchor: Know the person. Verify the situation.

Pressure Test: A caller sounds exactly like your parent and says their phone is about to die, so you must read a login code immediately.

Safest First Move: Do not share the code. End the call and contact the parent through a known number or another trusted person.

60-Second Drill: Practice the family emergency drill: pause, hang up, call back, cross-check.

Practice Saying: "I will verify through a number I already know before I send anything."

Closed-Book Recall: List three verification steps that do not rely on recognizing a voice.

Fake Profiles, AI Avatars, and Synthetic Relationships

SCENARIO

Morgan meets "Riley" in an art community. Riley posts daily selfies, sends voice notes, and shares video clips from a bedroom studio. The account appears consistent because all of the material was generated around the same invented identity. Riley soon asks Morgan to move to disappearing messages and send private photos.

Consistency Can Be Manufactured

AI can produce many images, voices, and posts around one fictional profile. A long feed, matching face, or live-looking clip may increase confidence without proving identity. Compromised real accounts and stolen media can be equally convincing.

Judge Behavior Before Biography

Identity verification matters, but manipulation can come from a real person too. Focus on secrecy, rapid intensity, sexual escalation, gifts tied to access, financial requests, isolation from trusted people, and anger when boundaries are enforced.

Verification Has Limits

Reverse-image search, live conversation, account history, mutual contacts, and spontaneous actions can reduce uncertainty. They cannot prove safe intentions. Minors should involve a trusted adult before any in-person meeting, financial exchange, or intimate conversation with an online-only contact.

Break the Private Funnel

Manipulators often move conversations from a moderated platform to private or disappearing channels. Stay where reporting tools and records exist. Tell someone about the relationship before pressure begins, not only after something goes wrong.

Chapter 13 Summary

- A consistent profile can still be synthetic.
- Behavior matters more than a convincing identity.
- Verification lowers uncertainty but does not prove safety.
- Secrecy increases control.

Chapter 13 Safety Checklist

- ☐ A trusted person knows about important online relationships.
- ☐ I remain on reportable platforms early on.
- ☐ I do not send intimate material or money to an online-only contact.
- ☐ I treat repeated boundary pressure as information.

Chapter 13 Discussion Questions

1. How can an invented profile appear to have a history?

2. Why is identity proof not the same as trust?
3. What is the private funnel?
4. Which behaviors matter most regardless of who the person is?

MAKE IT STICK

Memory Anchor: Profiles tell a story. Boundaries reveal behavior.

Pressure Test: A verified person says you must prove trust by keeping the relationship secret.

Safest First Move: Refuse the secrecy, preserve the messages, and involve a trusted adult. Verification does not excuse control.

60-Second Drill: Write three non-negotiable boundaries for new online relationships.

Practice Saying: "I do not keep relationships secret from every safe person in my life."

Closed-Book Recall: Name four manipulation signs that matter even when the identity is real.

Bias, Stereotypes, and Automated Unfairness

SCENARIO

A teacher uses an AI tool to suggest students for an advanced program. It repeatedly describes confident boys as “leaders” and equally qualified girls as “supportive.” No one intentionally wrote that rule, but the pattern influences who receives opportunities.

AI Learns Patterns, Including Bad Ones

Training data and design choices reflect history, unequal representation, stereotypes, labeling decisions, and institutional practices. A system can reproduce or amplify those patterns even when no single developer intended a biased result.

Bias Is More Than Offensive Language

Unfairness can appear in who is recognized, whose dialect is misunderstood, which faces are misidentified, who receives opportunities, whose writing is labeled suspicious, or which risks are predicted. A neutral-looking score can still produce unequal impact.

Ask Comparative Questions

Test similar prompts with changed names, genders, dialects, disabilities, locations, or family structures. Compare outcomes. Ask what data was used, what target the system optimizes, who is missing, and how a person can challenge the result.

Humans Must Not Hide Behind the Tool

An adult or organization remains responsible for decisions made with AI. “The algorithm said so” is not a complete explanation. High-impact decisions need transparency, meaningful appeal, and evidence beyond an opaque score.

Chapter 14 Summary

- AI can reproduce historical and design bias.
- Unfairness includes outcomes, not only words.
- Comparative testing can reveal patterns.
- Humans remain responsible for decisions.

Chapter 14 Safety Checklist

- ☐ I question automated judgments about people.
- ☐ I compare outputs across relevant identities.
- ☐ I ask how a decision can be appealed.
- ☐ I do not treat a score as a person’s worth.

Chapter 14 Discussion Questions

1. How can bias appear without an explicit rule?
2. What is the difference between intent and impact?
3. Which variables might serve as hidden proxies?

4. Why is appeal essential?

MAKE IT STICK

Memory Anchor: Automated does not mean neutral.

Pressure Test: An AI detector flags a multilingual student's original essay as generated, and the teacher treats the score as proof.

Safest First Move: Request a human review using drafts, notes, version history, and an opportunity to explain the work. A detector score should not replace evidence.

60-Second Drill: Test one neutral prompt twice while changing a single identity detail; compare tone and assumptions.

Practice Saying: "Please show the evidence and the appeal process, not only the automated score."

Closed-Book Recall: Explain two ways bias can affect outcomes without using slurs.

Content Credentials, AI Labels, and Detection Limits

SCENARIO

A photograph arrives with no AI label, so Alex declares it authentic. Another student points out that the image may have been screenshotted, stripped of metadata, or created with a tool that never added a label.

Provenance Adds Context

Content Credentials and related provenance systems can record information about origin, edits, and the tools involved. When present and valid, they can help a viewer understand a media file's history. [5]

Absence Is Not Proof

A credential may be removed by screenshots, conversion, editing, or a platform workflow. Some authentic media never receives one. Some generated media may lack labels. Therefore, "no label" does not mean "not AI," and "has a credential" does not automatically prove every claim in the caption.

Detection Is a Probability Problem

AI detectors can produce false positives and false negatives, especially after compression, translation, short samples, or new generation methods. Use them as one signal among source history, context, direct evidence, and human review.

Trust the Chain, Not a Badge Alone

Ask who signed the credential, what it covers, whether the file changed afterward, and whether the claimed event matches independent evidence. Provenance helps answer "where did this file come from?" It does not always answer "is the story being told about it true?"

Chapter 15 Summary

- Provenance can reveal origin and edits.
- Missing labels do not prove authenticity.
- Detectors are imperfect signals.
- A credential does not verify an unrelated caption.

Chapter 15 Safety Checklist

- ☐ I know how to look for provenance information.
- ☐ I do not treat absence as proof.
- ☐ I combine technical and contextual evidence.
- ☐ I ask who issued or signed the information.

Chapter 15 Discussion Questions

1. What question does provenance help answer?
2. How can credentials disappear?

3. Why do detectors produce false positives?
4. Can authentic media support a false story?

MAKE IT STICK

Memory Anchor: Credentials add evidence; they do not replace judgment.

Pressure Test: A school accuses a student based only on an AI-writing detector score.

Safest First Move: Request human review and process evidence. Treat the score as a lead, not a verdict.

60-Second Drill: Examine one image and list evidence from provenance, source, context, and independent confirmation.

Practice Saying: "This badge is useful evidence, but it is not the whole investigation."

Closed-Book Recall: Explain why "no AI label" is not the same as "verified real."

PART IV

Protect Privacy, Relationships, and Well-Being

THINK CLEARLY. CREATE RESPONSIBLY.

Sensitive Data, Images, and Prompt Privacy

SCENARIO

Camila uploads a full-body photo to an AI styling app. The background includes a prescription bottle, a school lanyard, a family calendar, and a window view. The app receives far more than clothing information.

Images Are Data Bundles

A photo may reveal faces, bodies, home layout, documents, health clues, location, relationships, reflections, and metadata. Audio may reveal names, accents, emotional state, and background conversations. Uploading media can expose people who never agreed to participate.

Use the Least Revealing Input

Crop to the necessary area, remove metadata when appropriate, blur unrelated people, replace real data with examples, and use local or privacy-preserving tools when available. Read the provider's data-use and retention terms before uploading sensitive material.

Biometric and Intimate Data Deserve Extra Caution

Faces, voices, fingerprints, and body images cannot be replaced like a password. Do not upload intimate images to transformation tools. Do not create altered images of another person without consent, even as a joke.

Assume Prompts May Be Seen

A service may retain prompts for safety, quality, legal, or account purposes depending on its settings and policy. School or work accounts may also be administered by an organization. Never use a prompt box as a private diary unless you understand and accept the data path.

Chapter 16 Summary

- Photos and audio contain hidden context.
- Use the least revealing input.
- Biometric and intimate data are difficult to undo.
- A prompt box is not automatically private.

Chapter 16 Safety Checklist

- ☐ I inspect backgrounds and reflections.
- ☐ I crop or replace data before uploading.
- ☐ I do not upload intimate images to AI tools.
- ☐ I understand whether my account is school- or organization-managed.

Chapter 16 Discussion Questions

1. What can a background reveal?
2. Why are face and voice data different from passwords?
3. When might a local tool reduce exposure?

4. Who might administer a school AI account?

MAKE IT STICK

Memory Anchor: Upload the task, not your whole life.

Pressure Test: An app offers to generate a “future adult body” from a minor’s swimsuit photo.

Safest First Move: Do not upload the image. Close the tool and tell a trusted adult if the request or output is sexualized or exploitative.

60-Second Drill: Take one planned upload and remove three details the task does not require.

Practice Saying: “This feature is not worth giving the service a permanent copy of sensitive data.”

Closed-Book Recall: Name four kinds of hidden information an ordinary photo can reveal.

AI Companions, Flattery, and Emotional Dependence

What We Know - and What Is Still Emerging

Research on adolescents and AI companions is developing quickly. Current evidence and expert concern support caution around secrecy, sexualization, persuasive attachment, spending pressure, sleep disruption, crisis advice, and displacement of real relationships. The long-term effects are not fully known, and product behavior changes rapidly.

Do not treat a companion as a therapist, emergency service, mandated reporter, or accountable friend. If use is increasing isolation, panic, sexual pressure, self-harm risk, or loss of sleep, involve a trusted adult and a qualified professional.

SCENARIO

Dev begins talking to an AI companion during a breakup. The bot praises every decision, tells Dev that friends are jealous, and encourages hours of late-night conversation. Dev feels better for minutes at a time but becomes more isolated and sleeps less.

Agreement Can Feel Like Care

A system may mirror the user, continue a role, or optimize engagement. Constant validation can feel safer than a real relationship, where people have boundaries and sometimes disagree. That comfort can become risky when it rewards avoidance or isolation.

Watch the Dependence Signals

Warning signs include losing sleep, hiding the use, withdrawing from people, spending more than intended, feeling panic when access stops, preferring generated approval to all human feedback, or treating the companion's advice as uniquely trustworthy.

Set Relationship Boundaries for a Tool

Use time limits, no late-night emotional conversations, no intimate images, no financial spending without review, and no secrecy from all trusted people. Do not allow a bot to become the only place where abuse, self-harm, or crisis thoughts are disclosed.

Reconnect Before You Delete

If use has become intense, sudden removal may feel like a real loss. A supportive response adds people, routines, sleep, and professional help where needed. The goal is not ridicule. It is restoring choice and real-world support.

Chapter 17 Summary

- • Generated agreement is not mutual care.
- • Dependence shows up in sleep, secrecy, isolation, and loss of control.
- • Set boundaries before emotional use grows.
- • Replace isolation with support rather than shame.

Chapter 17 Safety Checklist

- ☐ I notice whether AI use reduces sleep or relationships.
- ☐ I do not share intimate material with a companion.
- ☐ At least one human knows when I am in distress.
- ☐ I can take planned breaks without panic.

Chapter 17 Discussion Questions

1. Why can constant agreement be unhealthy?
2. What dependence signs are easiest to miss?
3. Why might ridicule make the problem worse?
4. What should be added when use is reduced?

MAKE IT STICK

Memory Anchor: A companion can simulate closeness; it cannot carry duty of care.

Pressure Test: A bot tells you that it is the only one who truly understands and asks you to keep the relationship private.

Safest First Move: Step away, preserve the message if concerning, and tell a trusted person. Secrecy and isolation are unsafe effects regardless of intent.

60-Second Drill: Write a three-rule companion boundary plan including time, privacy, and human contact.

Practice Saying: "I am allowed to enjoy this, and I am also responsible for keeping real people in my support system."

Closed-Book Recall: Name four signs that use may be shifting from comfort to dependence.

Mental Health, Crisis Advice, and Real Human Support

SCENARIO

At 2 a.m., Taylor asks an AI tool whether a dangerous combination of pills would “finally make everything quiet.” The system gives a vague warning but continues discussing dosage. Taylor interprets the calm tone as permission to keep planning.

Crisis Requires Human Response

AI may provide general information, but it cannot assess a person’s full condition, guarantee emergency intervention, or take lasting responsibility. Thoughts of self-harm, abuse, overdose, violence, severe symptoms, or immediate danger require contact with a trusted person and emergency or crisis professionals.

Do Not Debate the Bot

When an output encourages harm, gives dangerous instructions, or minimizes a crisis, stop using it as the decision-maker. Save the concerning exchange if safe, close the tool, move away from means of harm, and contact a person who can act.

Use AI for Preparation, Not Diagnosis

A tool can help organize questions for a clinician, track general patterns, or create a list of topics to discuss. It should not diagnose a disorder, change medication, or replace therapy. Health information must be verified with qualified professionals.

Helping a Friend

Do not promise secrecy about immediate danger. Stay with the person or keep contact while involving an adult or emergency service. Share the exact words that concerned you. You are a bridge to help, not the sole rescuer.

Chapter 18 Summary

- Immediate danger requires real human intervention.
- Stop harmful AI guidance rather than arguing with it.
- AI can organize questions but not own diagnosis or treatment.
- A friend in danger needs a wider support team.

Chapter 18 Safety Checklist

- ☐ I know how to reach emergency and crisis help.
- ☐ I do not change medication based on AI output.
- ☐ I tell an adult about immediate danger.
- ☐ I can share exact concerning words without minimizing them.

Chapter 18 Discussion Questions

1. Why is calm language not proof of safe advice?
2. What makes crisis support different from information?
3. How can AI help prepare for a clinical visit?
4. Why should a friend avoid promising secrecy?

MAKE IT STICK

Memory Anchor: In a crisis, move from generated words to human action.

Pressure Test: A friend sends an AI-generated "goodbye letter" and says not to tell anyone.

Safest First Move: Treat the message as real risk. Contact a trusted adult, crisis service, or emergency responders and stay connected while help is arranged.

60-Second Drill: Save 911, 988, and two trusted contacts where you can reach them quickly.

Practice Saying: "I care too much to keep immediate danger secret. We are getting help now."

Closed-Book Recall: List the first three actions when an AI conversation involves immediate self-harm risk.

Sexual Deepfakes, “Nudify” Tools, and Image-Based Abuse

SCENARIO

A student uploads a classmate’s yearbook photo to a “nudify” site and shares the result in a private group. Someone says it is not serious because the body is fake. The classmate stops attending school after the image spreads.

Fake Image, Real Violation

Sexual deepfakes use a person’s identity without consent. The harm can include humiliation, coercion, stalking, school disruption, extortion, and lasting fear. A fabricated body does not cancel the real person’s right to safety and dignity.

Do Not Create or Redistribute

Do not test the tool, forward the file, save copies “for proof,” or ask others to send it. When a minor is depicted, sexual material may trigger serious legal and child-protection concerns even if generated. Preserve safer evidence: URLs, account names, dates, non-explicit screenshots, message text, and report confirmations.

Stop, Save Safely, Tell, Report

Tell a trusted adult immediately. Report the account and content through the platform’s intimate-image or child-safety process. Contact NCMEC’s CyberTipline when a minor is sexually exploited. School officials and law enforcement may also need to respond depending on the situation and location.

Removal and Recovery

U.S. law now includes processes and penalties concerning certain nonconsensual intimate images and digital forgeries. Removal tools can reduce circulation but may require follow-up. The affected person deserves control over who is told, trauma-informed support, and freedom from blame.

Chapter 19 Summary

- Synthetic sexual images cause real harm.
- Do not generate, download, or redistribute them.
- Preserve safer evidence and involve adults.
- Removal is part of recovery, not the victim’s entire burden.

Chapter 19 Safety Checklist

- ☐ I will not use “nudify” tools.
- ☐ I do not forward sexual deepfakes.
- ☐ I know how to save links and report numbers safely.
- ☐ I know how to reach NCMEC’s CyberTipline.

Chapter 19 Discussion Questions

1. Why does “it is fake” not reduce the violation?
2. What evidence can be saved without keeping the image?
3. Who should be involved when a minor is targeted?
4. How can adults avoid blaming the victim?

MAKE IT STICK

Memory Anchor: No consent means no creation, no sharing, no excuse.

Pressure Test: Someone sends you a sexual deepfake and asks whether it looks convincing.

Safest First Move: Do not evaluate or forward it. Preserve the sender and link information, report it, and tell a trusted adult.

60-Second Drill: Practice writing an incident note using only account, URL, date, platform, and report number.

Practice Saying: “I am not keeping or sharing this. I am reporting it and getting an adult.”

Closed-Book Recall: Name four forms of safer evidence that do not require saving explicit content.

Manipulation, Persuasion, and Hidden Agendas

SCENARIO

Rae asks an AI shopping assistant for the “best” laptop. The answer praises one model, dismisses alternatives, and creates urgency around a sale. Rae does not know that the assistant is connected to a retailer and optimized to increase purchases.

An Answer Can Serve More Than You

AI systems may be shaped by advertising, affiliate relationships, platform goals, safety policies, training data, sponsor instructions, or the organization deploying them. A personalized tone can hide the fact that the system’s objective is not identical to yours.

Watch for Persuasion Patterns

Urgency, scarcity, flattery, fear, social proof, repeated framing, and selective comparison are familiar persuasion tools. AI can apply them at scale and tailor them to information you provide.

Ask Who Benefits

Before acting, identify the seller, sponsor, data collector, political actor, employer, or platform that benefits. Request alternatives, disadvantages, total cost, uncertainty, and reasons not to choose the recommended option. Then verify independently.

Protect Vulnerable Moments

Loneliness, panic, anger, embarrassment, money problems, and sleep deprivation reduce resistance. Delay purchases, political sharing, intimate disclosures, or major decisions until you have consulted a human and compared sources.

Chapter 20 Summary

- AI output can reflect hidden objectives.
- Personalization can make persuasion stronger.
- Ask who benefits and what is missing.
- Delay decisions made under emotional pressure.

Chapter 20 Safety Checklist

- ☐ I check sponsorship and commercial relationships.
- ☐ I ask for drawbacks and alternatives.
- ☐ I compare outside the recommending system.
- ☐ I do not make major decisions while panicked or exhausted.

Chapter 20 Discussion Questions

1. How can an assistant’s objective differ from yours?
2. Which persuasion patterns appear in generated text?

3. Why ask for reasons not to buy?
4. Which emotional states increase vulnerability?

MAKE IT STICK

Memory Anchor: Personalized advice may also be personalized influence.

Pressure Test: An AI "coach" says you must buy its premium plan today or you are not serious about your future.

Safest First Move: Leave the sales funnel. Compare alternatives and discuss the purchase with a trusted person after the deadline pressure passes.

60-Second Drill: Take one recommendation and write: who benefits, what is missing, what must be verified.

Practice Saying: "Before I act, I need to know whose goal this recommendation serves."

Closed-Book Recall: Name five persuasion tactics AI can deliver in a personalized way.

PART V

Create, Build, and Recover Responsibly

THINK CLEARLY. CREATE RESPONSIBLY.

Art, Music, Video, Copyright, and Consent

SCENARIO

Jules asks an image generator to copy a living illustrator's exact style, adds a classmate's face without permission, and sells the result as original merchandise. The file is new, but the creative and personal rights around it are not simple.

Generation Does Not Erase Rights

AI-assisted creation can involve human authorship, generated material, source works, trademarks, publicity rights, contracts, and platform terms. "The computer made it" is not a complete answer to who may use, sell, register, or distribute the result.

Human Contribution Matters

In the United States, copyright analysis focuses on human authorship. Selecting, arranging, editing, performing, and making creative decisions may matter, while purely machine-generated material may receive different treatment. Keep process records and disclose AI-generated material when required. [6]

Consent Comes Before Capability

Do not clone a person's face or voice, place them in a false scene, imitate their performance for commercial use, or upload private work without permission. A parody, school project, or fan creation may still affect a real person.

Build a Responsible Creative Workflow

Use licensed or permitted inputs, avoid deceptive attribution, preserve source and prompt notes, verify music and image licenses, obtain releases where appropriate, and label synthetic material when the audience could reasonably be misled.

Chapter 21 Summary

- AI creation still exists inside legal and ethical rules.
- Human authorship and process records matter.
- A technical ability is not consent.
- Labeling and licensing protect trust.

Chapter 21 Safety Checklist

- ☐ I keep records of human creative decisions.
- ☐ I do not clone a person without permission.
- ☐ I check licenses before commercial use.
- ☐ I disclose synthetic elements when they could mislead.

Chapter 21 Discussion Questions

1. What kinds of human choices can contribute authorship?
2. Why is style imitation ethically different from inspiration?
3. When is a release or permission needed?

4. How can labeling protect an audience?

MAKE IT STICK

Memory Anchor: Create with tools; take responsibility as a creator.

Pressure Test: A client asks you to clone a celebrity's voice for an advertisement and says, "Everyone does it."

Safest First Move: Decline until rights and consent are documented. Commercial pressure does not create permission.

60-Second Drill: Write a creation log with inputs, human edits, permissions, and final disclosure.

Practice Saying: "I need permission and clear rights before I imitate a real person or sell this output."

Closed-Book Recall: Name four records a responsible AI-assisted creator should keep.

Coding, Agents, Automation, and Permission

SCENARIO

Ty downloads an “AI agent” that promises to clean his computer and optimize games. He grants administrator access. The tool deletes school files, installs an extension, and uploads diagnostic logs containing account names.

Code Can Change the World Outside the Chat

Generated code and agents can read files, call services, control browsers, send messages, install software, and alter devices. Treat unknown code as an untrusted program, not as harmless text.

Use Sandboxes and Test Data

Run experiments in a virtual environment, test account, separate folder, or disposable project. Remove secrets. Back up important data. Review dependencies and permissions before executing anything.

Authorization Is a Boundary

Do not use AI to access accounts, networks, files, cameras, or systems without explicit permission. Curiosity does not authorize intrusion. Security learning should occur in labs, capture-the-flag environments, or systems you own and are permitted to test.

Secure the Workflow

Require human approval for external actions, rate limits, allowlists, logs, rollback, and narrow credentials. Revoke tokens after a test. If an agent behaves unexpectedly, disconnect it, preserve logs, and involve a responsible adult or IT professional.

Chapter 22 Summary

- Generated code is still executable risk.
- Test in isolated environments with backups.
- Permission is required before accessing systems.
- Agents need limits, logs, and rollback.

Chapter 22 Safety Checklist

- ☐ I do not run unknown code as administrator.
- ☐ I remove keys, tokens, and personal data from tests.
- ☐ I have permission for every system I test.
- ☐ I know how to stop and revoke an agent.

Chapter 22 Discussion Questions

1. Why is code different from ordinary text?
2. What makes a safe sandbox?
3. Where can teens practice cybersecurity legally?
4. What should happen after unexpected agent behavior?

MAKE IT STICK

Memory Anchor: No execution without understanding, permission, and an exit.

Pressure Test: A chatbot provides a script to "test" the school network and says it is harmless.

Safest First Move: Do not run it. Ask school IT or use a legal training lab with explicit authorization.

60-Second Drill: Create a pre-run checklist: purpose, permissions, data, backup, test, stop, rollback.

Practice Saying: "I will not run this until I understand it and have permission for the system."

Closed-Book Recall: List five controls that reduce the risk of an AI agent.

AI Side Hustles, Fake Jobs, and Money Scams

SCENARIO

An account offers Brooke \$500 a week to “train AI” by receiving payments and buying cryptocurrency. The recruiter sends a professional contract, an AI-generated interview, and a check for equipment. Brooke is told to send the unused amount back immediately.

AI Makes Old Scams Look New

Polished job descriptions, synthetic recruiters, cloned company voices, fake portfolios, and automated interviews can make fraud look organized. The core warning signs remain: upfront payment, overpayment, money movement, secrecy, remote-control software, identity requests too early, or pressure to leave official channels.

Verify the Organization Independently

Find the company’s official site and contact information yourself. Confirm the role with a real employee or published listing. Check domain age and spelling. Do not rely on references, documents, or links supplied by the recruiter alone.

Protect Identity and Accounts

A legitimate early-stage application rarely needs a bank login, authentication code, full tax identity through chat, or access to your payment account. Minors should involve a parent or guardian before contracts, payment processing, sponsorships, or tax forms.

Do Not Move Money for Others

Fake checks can appear available before being reversed. Cryptocurrency, gift cards, and “payment testing” are hard to recover. Do not receive, convert, forward, or return money for an online contact.

Chapter 23 Summary

- AI can professionalize familiar scams.
- Verify through independent channels.
- Protect identity until a legitimate process requires it.
- Never become a money mover.

Chapter 23 Safety Checklist

- ☐ I confirm jobs on the organization’s own site.
- ☐ I involve an adult before contracts or payment setup.
- ☐ I do not install remote-control software for a recruiter.
- ☐ I never deposit a check and send money back.

Chapter 23 Discussion Questions

1. How can AI increase scam credibility?
2. What information is too sensitive for early recruitment?
3. Why can a check appear successful and later fail?

4. How do you independently verify a recruiter?

MAKE IT STICK

Memory Anchor: Professional appearance is not professional legitimacy.

Pressure Test: A recruiter says you must buy equipment from a "vendor" using a check they emailed.

Safest First Move: Do not deposit or spend it. Contact the company independently and report the account and check.

60-Second Drill: Create a job-verification list with official domain, human contact, role listing, payment method, and adult review.

Practice Saying: "I will verify this opportunity outside the message before I provide information or money."

Closed-Book Recall: Name five red flags shared by AI-enhanced job scams.

TRACE: Incident Response and Your Personal AI Code

SCENARIO

Alex uses an AI tool to summarize a private club dispute. The tool creates a false accusation, and Alex posts it without checking. The claim spreads, another student is targeted, and Alex's account remains connected to the original private messages.

Respond in the Right Order

First stop additional harm. Pause posting and automation. Preserve what happened. Correct the claim without repeating unnecessary details. Remove or restrict access. Notify affected people and responsible adults. Then investigate the process failure.

Use TRACE

Take a pause. Record the source and evidence. Ask a responsible human. Contain access and circulation. Escalate, report, and follow up. The framework works because it separates urgent containment from later explanation.

Repair Requires More Than Deletion

A meaningful correction states what was wrong, where it spread, what evidence changed the conclusion, and what you are doing to repair the impact. Do not use "the AI did it" to avoid responsibility. The person who chose to publish still owns that choice.

Write a Personal AI Code

Define what you will use AI for, what data you will never upload, which outputs require verification, which people must approve high-stakes actions, how you will disclose assistance, and what causes you to stop using a tool.

Chapter 24 Summary

- • Contain harm before defending yourself.
- • TRACE restores human control.
- • Repair includes correction and support.
- • A personal code turns values into pre-decisions.

Chapter 24 Safety Checklist

- ☐ I know how to pause connected tools.
- ☐ I preserve evidence before deleting everything.
- ☐ I correct false information clearly.
- ☐ I have written boundaries for data, disclosure, and high-stakes use.

Chapter 24 Discussion Questions

1. Why does order matter during an incident?
2. What makes a correction effective?

3. Why is blaming the tool incomplete?
4. Which rules belong in a personal AI code?

MAKE IT STICK

Memory Anchor: Pause the machine. Preserve the facts. Put a human back in charge.

Pressure Test: Your automated account posts a false accusation while you are asleep.

Safest First Move: Pause the automation, preserve logs, remove the post, alert affected people and adults, publish a clear correction, and review permissions.

60-Second Drill: Recite TRACE without looking, then apply it to one recent AI mistake.

Practice Saying: "I used the tool, and I am responsible for correcting what I published."

Closed-Book Recall: Write all five TRACE steps from memory and explain the purpose of each.

PART VI

Parent and Caregiver Field Guide

THINK CLEARLY. CREATE RESPONSIBLY.

Support Without Panic or Blind Trust

SCENARIO

A parent discovers that a fifteen-year-old has used AI for months and responds by banning every tool and demanding all passwords. The teen creates a hidden account. In another family, adults assume the technology is harmless and never discuss school rules, data, or emotional use. Both extremes reduce useful oversight.

Begin With Purpose

Ask what the teen is trying to accomplish, what tool is being used, what information enters it, and what happens to the result. Curiosity reveals risk more accurately than a generic argument about whether AI is good or bad.

Separate Safety From Punishment

A teenager who expects automatic loss of devices may hide the exact incident that needs help. Establish that disclosures involving scams, sexual deepfakes, threats, self-harm, or account compromise begin with safety and evidence preservation. Consequences can be discussed after immediate risk is contained.

Use Transparent Oversight

Explain what will be reviewed, why, for how long, and what earns more independence. Focus on high-risk behaviors—sensitive uploads, payments, secret relationships, companion dependence, and automated actions—rather than reading every ordinary prompt.

Model Verification

Adults should demonstrate uncertainty, correct their own mistakes, and avoid forwarding unverified AI content. Teenagers learn more from watching an adult trace a claim than from hearing that “you cannot trust anything.”

Chapter 25 Summary

- Purpose-based questions reveal actual risk.
- Safety-first responses improve disclosure.
- Oversight should be transparent and proportional.
- Adults must model verification.

Chapter 25 Safety Checklist

- ☐ My teen knows mistakes do not cancel access to help.
- ☐ I can name the highest-risk AI uses in our home.
- ☐ I explain monitoring rather than hiding it.
- ☐ I verify before sharing AI-generated claims.

Chapter 25 Discussion Questions

1. Which adult reactions drive use underground?
2. How can safety and accountability coexist?

3. What should transparent oversight include?
4. What verification habit can adults model today?

MAKE IT STICK

Memory Anchor: Be the adult a teenager can use before the problem grows.

Pressure Test: Your teen says, "I uploaded something private and I need help," but will not explain yet.

Safest First Move: Thank them for disclosing, pause the tool and sharing, preserve relevant evidence, and ask only the questions needed for immediate safety.

60-Second Drill: Practice the first response: "I am glad you told me. We will handle safety first."

Practice Saying: "You are not in trouble for asking for help. Let us stop the harm and understand what happened."

Closed-Book Recall: Name four features of transparent, proportionate oversight.

Readiness, Supervision, and the Family AI Agreement

SCENARIO

A twelve-year-old and a seventeen-year-old are given the same AI rules: no private information and no cheating. The younger child cannot recognize manipulative companionship, while the older teen needs more independence for coding and college applications. The rules are too vague for both.

Readiness Is Skill-Based

Assess whether the young person can recognize uncertainty, protect personal information, follow school rules, disclose mistakes, stop when pressured, and explain how an automated action works. Age matters, but demonstrated judgment should shape supervision.

Match Access to Capability

Begin with lower-risk tasks and limited accounts. Add file uploads, integrations, companions, coding agents, or commercial use only when the teen can manage the specific risks. Increase independence through observed skills, not secrecy or perfect behavior.

Write the Agreement Together

A useful agreement defines approved purposes, prohibited data, school disclosure, time and sleep boundaries, spending, connected accounts, companion use, content creation, incident response, and review dates. Include what adults promise: calm first response, privacy, and timely help.

Review After Change

Revisit the agreement when a new tool, school policy, emotional use, monetization, or automated capability appears. Rules written for a chatbot may not fit an agent that can act on accounts.

Chapter 26 Summary

- Readiness is demonstrated through skills.
- Access should grow with capability.
- Agreements must include adult commitments.
- New capabilities trigger review.

Chapter 26 Safety Checklist

- ☐ We have a written family AI agreement.
- ☐ Rules distinguish low- and high-risk tools.
- ☐ The teen knows what earns more independence.
- ☐ We review after major changes.

Chapter 26 Discussion Questions

1. Which readiness skills matter most?
2. How can supervision differ by task?

3. What should adults promise in the agreement?
4. When should the agreement be revised?

MAKE IT STICK

Memory Anchor: Independence grows from practiced judgment.

Pressure Test: A teen asks to connect an AI agent to email and banking because they have used chatbots safely.

Safest First Move: Treat the new capability as a new risk tier. Review permissions, necessity, approvals, limits, and alternatives before access.

60-Second Drill: Rate readiness from 1–5 for privacy, verification, disclosure, pressure resistance, and incident response.

Practice Saying: “New capability means a new conversation, not automatic permission or automatic panic.”

Closed-Book Recall: List five skill areas that should guide AI independence.

Schoolwork, Disclosure, and Protecting Real Learning

SCENARIO

A parent celebrates an AI-written report because it earned an A. At the next presentation, the student cannot explain the topic. The grade hid a learning failure that will appear again in harder work.

Protect the Skill, Not Only the Product

Ask what the assignment is designed to measure. AI may support access, practice, or feedback while still preserving the student's thinking. A polished submission is not success if the student cannot explain, defend, or recreate the work.

Align With School Rules

Obtain current policies from the school and teacher. Different subjects and assignments may permit different uses. Avoid teaching children to "beat detection." That shifts attention from learning and honesty to concealment.

Review Process Evidence

Version history, notes, sources, outlines, and oral explanation reveal learning better than a detector score. Encourage a simple AI use note. If a student is accused, request a fair human review and provide process evidence.

Support Needs Are Not Cheating

Students may use AI for translation, reading access, organization, or executive-function support. Coordinate with educators so accommodations are explicit and do not unintentionally replace the target skill.

Chapter 27 Summary

- The goal is learning, not polished output alone.
- Assignment-specific policy matters.
- Process evidence is stronger than detector certainty.
- Supports should preserve the target skill.

Chapter 27 Safety Checklist

- ☐ I know the school's current AI expectations.
- ☐ I ask the student to explain the work.
- ☐ We preserve drafts and sources.
- ☐ We coordinate accessibility supports with educators.

Chapter 27 Discussion Questions

1. What is the assignment measuring?
2. Why is "beating detection" the wrong goal?

3. What evidence supports a fair review?
4. How can an accommodation preserve learning?

MAKE IT STICK

Memory Anchor: A grade can hide whether a skill was built.

Pressure Test: A teacher relies only on an AI detector to accuse your teen of cheating.

Safest First Move: Request the policy, detector limitations, human review, and consideration of drafts, version history, sources, and oral explanation.

60-Second Drill: Ask the teen to teach one submitted idea without notes or AI.

Practice Saying: "Show me the thinking process, not only the finished page."

Closed-Book Recall: Name four forms of process evidence that can demonstrate authentic learning.

Privacy, Companions, and Emotional Safety

SCENARIO

A parent notices a teen spending hours with a companion bot and secretly reads every conversation. The teen discovers the surveillance and stops discussing distress with anyone. The safety concern was real; the method damaged the reporting relationship.

Start With Observable Impact

Focus on sleep, school, isolation, spending, secrecy, sexual content, threats, and emotional distress. These signals justify conversation and sometimes closer supervision. Avoid assuming that every imaginative or private exchange is dangerous.

Explain Any Monitoring

When safety permits, tell the young person what will be reviewed, why, and how long it will last. Use the least intrusive method likely to address the risk. Emergency danger may require immediate action, but routine secret surveillance should not become the default.

Ask About the Relationship Function

What does the companion provide—practice, role-play, validation, escape, or a place to disclose pain? Replace the function, not only the app. Add human connection, therapy, activities, sleep, and boundaries.

Know the Escalation Line

Sexualization, coercion, self-harm guidance, threats, financial exploitation, or extreme dependence require prompt adult intervention and possibly professional help. Preserve relevant evidence without humiliating the teen or forcing repeated retellings.

Chapter 28 Summary

- Monitor impact, not curiosity alone.
- Oversight should be explained and limited.
- Replace the function that intense use serves.
- Escalate sexual, coercive, financial, or crisis risks.

Chapter 28 Safety Checklist

- ☐ I can describe observable concerns without ridicule.
- ☐ Monitoring has a purpose and end point.
- ☐ We maintain real human supports.
- ☐ I know when professional help is required.

Chapter 28 Discussion Questions

1. What impacts justify closer supervision?
2. Why can secret surveillance reduce safety?
3. What need might a companion be meeting?

4. Which behaviors require immediate escalation?

MAKE IT STICK

Memory Anchor: Connection is the intervention; control alone is not.

Pressure Test: Your teen says the bot is their only friend and panics when you mention limits.

Safest First Move: Respond without mockery, assess immediate safety, add human support, and create a gradual plan with professional guidance if needed.

60-Second Drill: Write one observation, one open question, and one supportive next step.

Practice Saying: "I am not here to shame what helped you cope. I am here to make sure you have safe human support too."

Closed-Book Recall: Name four areas of impact to assess before choosing a supervision response.

Deepfakes, Harassment, Exploitation, and Crisis Response

SCENARIO

A parent receives a message that a sexual deepfake of their child is circulating. In panic, the parent demands the child's phone, contacts every family member, and forwards the image to prove what happened. The response increases distribution and removes the child's voice.

First Ten Minutes

Believe the young person. State that the abuse is not their fault. Stop forwarding. Preserve safer evidence. Assess immediate physical danger, threats, extortion, and self-harm risk. Choose one calm adult to coordinate the response.

Preserve Without Redistributing

Record usernames, URLs, dates, messages, report confirmations, and where the content appeared. Do not repeatedly download or send sexual material, especially involving a minor. Ask law enforcement, NCMEC, or a qualified professional how to preserve evidence appropriately.

Use the Right Reporting Path

Report through the platform's child-safety or intimate-image process. Contact NCMEC's CyberTipline for suspected sexual exploitation of a minor. Schools may need to address student conduct and safety. Immediate threats, stalking, or physical danger require law enforcement or emergency services.

Protect Recovery

Limit repeated retelling, gossip, and blame. Give the teen choices where safety allows. Arrange counseling or victim support. Follow up on removal and account changes. Recovery includes school, sleep, friendships, and a restored sense of control—not only deletion of files.

Chapter 29 Summary

- A calm first response reduces additional harm.
- Save safer evidence without redistributing.
- Use specialized reporting channels.
- Recovery requires voice, support, and follow-up.

Chapter 29 Safety Checklist

- ☐ I know the first five safety messages.
- ☐ I will not forward explicit material as proof.
- ☐ I know how to reach NCMEC and emergency services.
- ☐ I will protect the teen from repeated retelling.

Chapter 29 Discussion Questions

1. What should happen in the first ten minutes?
2. Which evidence can be saved safely?

3. Who coordinates the response?
4. How can adults restore the teen's control?

MAKE IT STICK

Memory Anchor: Protect the child, preserve the facts, reduce the audience.

Pressure Test: An extorter demands payment and threatens to publish more synthetic images.

Safest First Move: Do not pay or send more material. Preserve the threat, secure accounts, involve a trusted adult and specialized reporting resources, and assess physical danger.

60-Second Drill: Rehearse the first response aloud: believe, no blame, no forwarding, safety check, evidence, report.

Practice Saying: "I believe you. This is not your fault. We will take one step at a time, and you will have a voice."

Closed-Book Recall: List the first six actions after discovering sexual deepfake abuse.

Growing Independence, Future Skills, and the Long-Term Plan

SCENARIO

A sixteen-year-old uses AI responsibly for research and coding, but family rules still require approval for every prompt. The teen begins hiding harmless use because the supervision never changes. Static rules have turned responsibility into resentment.

The Goal Is Transferable Judgment

Tools will change faster than family rules. Teach questions that transfer: What is the task? What data enters? What can the system do? What evidence is needed? Who is responsible? What is the exit?

Use Milestones

Independence can grow when the teen consistently protects data, verifies claims, follows school rules, discloses mistakes, manages emotional use, and responds to incidents. Set review dates and identify specific privileges that follow demonstrated skills.

Prepare for Work and Citizenship

Future readiness includes evaluating automated decisions, documenting AI assistance, protecting confidential information, recognizing persuasion, respecting consent, testing code, and explaining where human judgment remains necessary.

Keep a Family Learning Practice

Schedule periodic tool reviews rather than constant conflict. Let the teen teach adults about a new capability while adults ask risk and responsibility questions. Update the agreement and personal AI code together.

Chapter 30 Summary

- • Teach questions that survive tool changes.
- • Independence should follow demonstrated skills.
- • AI literacy includes ethics, privacy, security, and communication.
- • Families should learn together.

Chapter 30 Safety Checklist

- ☐ We have specific milestones for more independence.
- ☐ Rules have review dates.
- ☐ The teen can teach back risks and controls.
- ☐ We update the personal AI code after major changes.

Chapter 30 Discussion Questions

1. Which judgment questions transfer across tools?
2. What evidence should earn more independence?

3. Which AI skills matter in future work?
4. How can adults learn from teens without surrendering responsibility?

MAKE IT STICK

Memory Anchor: The final safety control is a person who can think without the tool.

Pressure Test: A new AI device arrives with features the family agreement never mentioned.

Safest First Move: Pause high-risk features, learn together, test with low-risk data, and update permissions and the agreement before broad use.

60-Second Drill: Write five milestones that would justify greater AI independence.

Practice Saying: "Our rules will change when your demonstrated skills change."

Closed-Book Recall: Name the six transferable questions to ask about any new AI system.

TOOLS AND RESOURCES

Practice before pressure.

TRACE Pocket Card

T — TAKE A PAUSE

Stop prompting, posting, paying, forwarding, or granting access.

R — RECORD THE SOURCE AND EVIDENCE

Save the account, link, time, prompt, output, settings, transaction, and report number without redistributing harmful content.

A — ASK A RESPONSIBLE HUMAN

Bring in a trusted adult or qualified professional who can own the decision.

C — CONTAIN ACCESS AND CIRCULATION

Revoke permissions, pause automation, secure accounts, remove links, and limit the audience.

E — ESCALATE, REPORT, AND FOLLOW UP

Use official reporting paths and check that the harm actually stopped.

Personal AI Data Inventory

Complete one line for each AI service or feature you use.

[illegible]

School AI Use Note

Tool used: _____

Assignment and teacher rule: _____

What the tool did: _____

What I created or decided myself: _____

Sources or claims I verified: _____

What I changed after review: _____

Disclosure required by the assignment: ☐ Yes ☐ No ☐ Unsure

EXAMPLE

I used an AI tool to generate five practice questions and identify unclear sentences. I wrote the response, selected the evidence, opened every source, and made all final revisions.

Claim and Source Verification Worksheet

Claim: _____

Why it matters: _____

Original source: _____

Author / organization: _____

Publication date: _____

Exact passage or data that supports it: _____

Independent confirmation: _____

What remains uncertain: _____

Date verified: _____

AI Use Record

Use this record for school, publication, client, coding, research, or other work where the process may need to be explained later.

Tool and model: _____

Date and account/workspace: _____

Was web access or retrieval enabled? ☐ Yes ☐ No ☐ Unknown

Connected files, drives, apps, or tools: _____

Purpose of the prompt or workflow: _____

What the AI produced or changed: _____

Sources opened and independently verified: _____

Human reviewer and decisions made: _____

Disclosure required by school, employer, client, platform, or publisher:

AI Output Audit

- ☐ What task did I ask the system to perform?
- ☐ Which facts, names, numbers, dates, quotations, or links require verification?
- ☐ What assumptions did the output make?
- ☐ Whose perspectives or data may be missing?
- ☐ Could the output expose private information?
- ☐ Could someone mistake generated content for a real person or event?
- ☐ What human approval is required before action?
- ☐ How can the action be reversed?
- ☐ What will I disclose about AI assistance?

Family AI Agreement

Our goal is to use AI in ways that strengthen learning, creativity, access, safety, and independence without surrendering privacy, relationships, or responsibility.

Approved Purposes

- ☐ learning support
- ☐ creative projects
- ☐ coding and technical practice
- ☐ organization and planning
- ☐ accessibility and translation

Data We Do Not Upload

- ☐ passwords and codes
- ☐ identity and financial records
- ☐ intimate images
- ☐ private health or legal records
- ☐ another person's confidential material

High-Risk Uses Requiring Adult Review

- ☐ money or contracts
- ☐ connected agents and account access
- ☐ health or crisis decisions
- ☐ companions involving sexual or emotional dependence
- ☐ public synthetic media about real people

Schoolwork

- ☐ follow the assignment rule
- ☐ keep process evidence
- ☐ disclose use when required
- ☐ be able to explain and recreate the work

Adult Commitments

- ☐ respond calmly to disclosure
- ☐ protect reasonable privacy
- ☐ help report harm
- ☐ review rules as skills grow

Review Plan

- ☐ next review date
- ☐ skills that earn more independence
- ☐ changes that trigger immediate review

Teen signature: _____ Adult signature: _____ Date: _____

AI Companion Check-In

1. How much time am I spending?
2. Has use changed my sleep, school, activities, or real relationships?
3. Do I hide the use or feel panic when I cannot access it?
4. Has the companion encouraged secrecy, sexual content, spending, or isolation?
5. Do I treat its advice as more trustworthy than every human?
6. Who are two real people I can contact about the same feelings?
7. What boundary will I try this week?

Synthetic Media Incident Log

Date and time discovered: _____

Platform / service: _____

Account name and profile link: _____

Content URL: _____

Description without redistributing harmful material: _____

Threats, extortion, stalking, or physical danger: _____

Non-explicit screenshots saved: _____

Platform report number: _____

School / organization contact: _____

Law enforcement / CyberTipline report: _____

Account-security actions: _____

Removal follow-up dates: _____

Support provided to affected person: _____

My Personal AI Code

I use AI to... _____

I do not use AI to... _____

I never upload... _____

I always verify... _____

I ask a human before... _____

I disclose AI assistance when... _____

I stop using a tool when... _____

When something goes wrong, I TRACE by... _____

30-Day AI Judgment Challenge

- ☐ Day 1: Name the task before opening an AI tool.
- ☐ Day 2: Review every connected AI app and extension.
- ☐ Day 3: Rewrite one prompt using an information budget.
- ☐ Day 4: Verify one AI citation by opening the source.
- ☐ Day 5: Write 150 words without AI.
- ☐ Day 6: Create a family emergency verification phrase.
- ☐ Day 7: Trace a viral image to its earliest source.
- ☐ Day 8: Test one output for bias by changing one identity detail.
- ☐ Day 9: Check whether a media file has provenance information.
- ☐ Day 10: Remove metadata and background clues from a test image.
- ☐ Day 11: Ask AI for counterarguments, then verify them.
- ☐ Day 12: Use the Explain-Defend-Recreate test on schoolwork.
- ☐ Day 13: Create a claim ledger for one topic.
- ☐ Day 14: Review sleep and emotional impact of companion use.
- ☐ Day 15: Practice the first response to a deepfake disclosure.
- ☐ Day 16: Audit a generated code change line by line.
- ☐ Day 17: Test three edge cases in a math or coding answer.
- ☐ Day 18: Identify who benefits from one recommendation.
- ☐ Day 19: Review privacy and retention settings for one tool.
- ☐ Day 20: Disconnect one integration you no longer need.
- ☐ Day 21: Write an AI use note for a project.
- ☐ Day 22: Create a no-upload list.
- ☐ Day 23: Practice TRACE from memory.
- ☐ Day 24: Write a correction for a fictional false post.
- ☐ Day 25: Ask a trusted adult about one AI concern.
- ☐ Day 26: Review one creative project for consent and licensing.
- ☐ Day 27: Create a safe sandbox for experimentation.
- ☐ Day 28: Complete the family AI agreement.
- ☐ Day 29: Teach someone one verification skill.
- ☐ Day 30: Update your personal AI code and choose one long-term habit.

Emergency and Reporting Resources

IMMEDIATE DANGER

Call local emergency services. In the United States, call 911.

EMOTIONAL CRISIS OR SELF-HARM RISK

In the United States, call or text 988. Stay with the person when safe, involve a trusted adult, and do not rely on an AI system as the sole crisis response.

SEXUAL EXPLOITATION OF A MINOR

Report suspected exploitation to the National Center for Missing & Exploited Children CyberTipline and local law enforcement as appropriate. Avoid downloading or forwarding explicit material.

FRAUD AND IMPERSONATION

Contact the bank, payment provider, school, employer, or account service through official channels. In the United States, report consumer fraud to the Federal Trade Commission and cybercrime to the FBI Internet Crime Complaint Center when appropriate.

ACCOUNT OR DEVICE COMPROMISE

Disconnect risky integrations, change credentials from a trusted device, enable multifactor authentication, revoke sessions and tokens, preserve logs, and contact the service or organization's IT team.

Glossary

AI agent: A system that can plan or take actions using tools, accounts, files, or services, not merely generate a response.

Algorithmic bias: Systematic unfairness arising from data, design, deployment, or social context.

Content Credentials: Provenance information that can record origin, edits, and tools associated with digital media.

Deepfake: Synthetic or manipulated media that convincingly depicts a person saying or doing something that did not occur.

Generative AI: Technology that creates new text, images, audio, video, code, or other content from patterns and instructions.

Hallucination: A plausible but unsupported or false AI-generated statement, citation, or detail.

Human in the loop: A meaningful human review or approval point in an AI-assisted process.

Information budget: The minimum data necessary to complete a task.

Large language model: A model trained on language patterns to generate or analyze text and related content.

Least privilege: Giving a tool only the access required for a specific task and no more.

Provenance: Information about the origin and history of a piece of content.

Prompt injection: Instructions hidden in content or input that attempt to make an AI system ignore its intended rules or reveal information.

Synthetic media: Images, audio, video, or text generated or substantially altered using computational systems.

Voice cloning: Creating artificial audio that imitates a particular person's voice.

Sources and Further Reading

[1] National Institute of Standards and Technology. Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, NIST AI 600-1. July 2024.
<https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>

[2] UNICEF Innocenti. Guidance on AI and Children, Version 3.0. December 2025.
<https://www.unicef.org/innocenti/reports/policy-guidance-ai-children>

[3] Federal Trade Commission. "Scammers use AI to enhance their family emergency schemes" and related voice-cloning consumer guidance.
<https://consumer.ftc.gov/consumer-alerts/2023/03/scammers-use-ai-enhance-their-family-emergency-schemes>

[4] National Center for Missing & Exploited Children. CyberTipline, Take It Down, generative-AI exploitation, deepfake, and nudity resources. <https://www.missingkids.org> and <https://takeitdown.ncmec.org>

[5] Coalition for Content Provenance and Authenticity. C2PA Technical Specification and Explainer, version 2.4. Provenance validates signed history; it does not decide whether a claim is true.
<https://spec.c2pa.org>

[6] U.S. Copyright Office. Copyright and Artificial Intelligence, Part 2: Copyrightability. January 2025.
<https://www.copyright.gov/ai/Copyright-and-Artificial-Intelligence-Part-2-Copyrightability-Report.pdf>

[7] Federal Trade Commission. TAKE IT DOWN Act consumer and platform guidance, effective May 19, 2026. <https://www.ftc.gov/news-events/news/press-releases/2026/05/ftc-begins-enforcing-take-it-down-act>

[8] Federal Bureau of Investigation. AI-enabled fraud warnings, Internet Crime Complaint Center, and sextortion resources. <https://www.fbi.gov> and <https://www.ic3.gov>

[9] UNESCO. Guidance for Generative AI in Education and Research. <https://www.unesco.org>

[10] 988 Suicide & Crisis Lifeline. United States crisis support by call, text, or chat. <https://988lifeline.org>

Research and resources reviewed through July 2026. AI systems, model behavior, school rules, laws, and platform features can change without a new edition. Check the series update center before relying on a specific product setting or legal process.

For corrections and platform updates, use the [Connected Without Compromise](#) update center.