

Wer hat das sortiert?

Woher Kategoriensysteme kommen — und wie man ein gewachsenes Vokabular gegen sie prüft

Lars Tiedemann

Dipl. Heilpädagoge (FH) · AssistUK

22. August 2026

Arbeitsfassung. Dieser Text ist der Schwestertext zu »Wo die Wörter wohnen«. Dort geht es darum, wie Menschen Kategorien bilden — Prototypen, Basisebene, Ereignisschemata. Hier geht es um die andere Hälfte derselben Frage: Woher die fertigen Kategoriensysteme kommen, die wir benutzen, wer sie gebaut hat und wofür, und wie man ein über Jahre gewachsenes Vokabular systematisch gegen sie prüft, ohne sich selbst zu belügen. Der Anlass war ein konkretes Audit unseres eigenen Bestands; die Zahlen in Abschnitt 9 und 10 sind gemessen, nicht geschätzt. Alle Quellen wurden am 22.08.2026 auf Existenz und Kernaussage geprüft; das Verzeichnis am Ende weist den Prüfstatus aus.

Inhalt

1. Eine Ordnung, die niemand geprüft hat
 2. Die älteste Antwort: das Wortfeld
 3. Prototypen in drei Sätzen
 4. Buck und die IDS: eine Feldliste für den Sprachvergleich
 5. WordNet: dieselbe Idee, maschinenlesbar
 6. Der Farbttest: warum »bunt« keine Farbe ist
 7. Die zweite Achse: Wortart statt Bedeutung
 8. Ein Prüfverfahren in drei Achsen
 9. Warum eine rohe Fehlerquote lügt
 10. Was dabei herauskam
 11. Wie wir es bei AssistUK halten
 12. Grenzen dieses Textes
 13. Fazit in fünf Sätzen
 14. Literatur (mit Prüfvermerk)
-

1. Eine Ordnung, die niemand geprüft hat

Jedes Symbolvokabular, das über Jahre wächst, sammelt seine Kategorien nebenbei ein. Eine Liste entsteht, weil jemand Wörter für ein Thema brauchte. Eine zweite wird abgespalten, weil die erste zu groß wurde. Eine dritte kommt aus einem Vorbild-System. Nach einigen Jahren steht ein Kategoriensystem da, das niemand als Ganzes entworfen hat — und das trotzdem täglich darüber entscheidet, ob ein Mensch sein Wort findet oder nicht.

Bei unserem eigenen Bestand waren es zuletzt 31 Hauptgruppen, 241 Listen und rund 8.600 Einträge. Ich konnte für jede einzelne Liste sagen, warum es sie gibt. Ich konnte nicht sagen, ob das Ganze eine Ordnung ist.

Das ist die Frage, um die es hier geht. Sie klingt akademisch und ist es nicht: Wer sein Wort nicht findet, sagt es nicht. Die Kategorie ist kein Etikett, sie ist der Suchweg.

Nun ließe sich ein solches System nach Gefühl prüfen — man geht die Listen durch und stutzt dort, wo etwas komisch aussieht. Das habe ich lange so gemacht, und es findet die groben Fälle. Es findet nicht die systematischen. Und es hat einen Fehler, den man sich selbst nur schwer nachweist: Man prüft die eigene Ordnung mit dem Kopf, der sie gebaut hat.

Was es dafür braucht, ist ein **Maßstab von außen**. Nicht, weil externe Systeme klüger wären, sondern weil sie unabhängig sind.

2. Die älteste Antwort: das Wortfeld

Der Gedanke, dass Wörter nicht einzeln im Kopf liegen, sondern in Feldern, ist fast hundert Jahre alt. Jost Trier hat ihn 1931 in seiner Untersuchung zum deutschen Wortschatz »im Sinnbezirk des Verstandes« ausformuliert: Ein Wort hat seine Bedeutung nicht für sich,

sondern durch die Nachbarn, gegen die es sich abgrenzt. Wer *klug* sagt, sagt damit auch: nicht *weise*, nicht *listig*, nicht *gebildet*. Ändert sich ein Nachbar, ändert sich das Wort mit.

John Lyons hat das später nüchterner gefasst und in eine Semantik eingebettet, die zwischen Sinnrelationen unterscheidet — Über- und Unterordnung, Gegensatz, Teil-Ganzes. Damit wurde aus einer sprachphilosophischen Intuition ein beschreibbares Werkzeug.

Für uns folgt daraus zweierlei. Erstens: Eine Kategorie ist kein Behälter, in den man Dinge legt, sondern ein **Nachbarschaftsverhältnis**. Das erklärt, warum es sich falsch anfühlt, wenn *hell* und *dunkel* auf einer Seite weit auseinanderstehen — sie definieren sich gegenseitig. Zweitens: Felder sind sprachabhängig. Das deutsche Feld um *Kopf* deckt sich nicht mit dem englischen um *head*. Wer zweisprachige Vokabulare baut, hat dieses Problem an jeder Zelle.

3. Prototypen in drei Sätzen

Die zweite große Antwort kommt aus der Psychologie und ist im Schwestertext ausführlich beschrieben; hier nur so viel, wie man zum Weiterlesen braucht.

Kategorien haben keine scharfen Grenzen und keine Liste notwendiger Merkmale. Sie haben ein **Zentrum** und einen unscharfen Rand: Ein Spatz ist ein besserer Vogel als ein Pinguin, ein Stuhl ein besseres Möbelstück als eine Stehlampe. Und es gibt eine **Basisebene**, auf der Menschen von selbst benennen — *Hund*, nicht *Säugetier* und nicht *Dackel*.

Wittgenstein hatte das Prinzip vorweggenommen, als er zeigte, dass sich die Dinge, die wir »Spiel« nennen, kein einziges gemeinsames Merkmal teilen, sondern ein Netz von Familienähnlichkeiten. Eleanor Rosch hat es experimentell belegbar gemacht. Für die Praxis heißt es: Wer eine Kategorie sucht, sucht ihr **typisches** Mitglied. Eine Kachel muss darum ihren Prototypen zeigen, nicht ihren Grenzfall.

4. Buck und die IDS: eine Feldliste für den Sprachvergleich

Wortfelder als Theorie helfen noch nicht beim Sortieren. Man braucht eine Liste. Und es gibt eine, die seit Jahrzehnten für genau den Zweck gebaut wird, den wir haben.

Carl Darling Buck veröffentlichte 1949 ein *Dictionary of Selected Synonyms in the Principal Indo-European Languages*. Sein Ziel war nicht Vollständigkeit, sondern Vergleichbarkeit: Er wollte wissen, wie verschiedene Sprachen dieselben Grundbedeutungen ausdrücken. Dafür ordnete er rund 1.200 Bedeutungen in **22 thematische Kapitel** — die physische Welt, Verwandtschaft, Tiere, der Körper, Essen und Trinken, Kleidung, das Haus, Landwirtschaft, grundlegende Handlungen, Bewegung, Besitz, räumliche Beziehungen, Menge, Zeit, Sinneswahrnehmung, Gefühle und Werte, Denken, Sprache, gesellschaftliche Beziehungen, Krieg und Jagd, Recht, Religion.

Diese Gliederung lebt weiter in der **Intercontinental Dictionary Series** (IDS), begründet von Mary Ritchie Key und heute von Bernard Comrie herausgegeben. Die IDS übernimmt Bucks thematischen Aufriss samt Nummerierung und erweitert ihn auf **1.310 Bedeutungen in 22 Kapiteln**; fehlt einer Sprache ein Wort, bleibt der Eintrag leer. Das Material steht offen zur Verfügung.

Warum das für uns taugt, ist kein Zufall, sondern Konstruktionsprinzip: Die Liste ist für **Grundwortschatz über Sprachgrenzen hinweg** gebaut. Genau das ist unsere Aufgabe — ein Vokabular, das in zweiunddreißig Sprachen dieselben Zellen trägt. Ein Kategoriensystem, das nur im Deutschen aufgeht, wäre für uns wertlos.

Zwei Einschränkungen sind ehrlicherweise zu nennen. Bucks Ordnung ist ein Erzeugnis ihrer Zeit und ihres Gegenstands: Ein Kapitel »Krieg und Jagd« steht in einem Kommunikationsvokabular für Menschen mit Behinderungen anders da als in einer indogermanistischen Wortstudie. Und die Liste kennt die moderne Lebenswelt nicht — Computer, Fernsehen, Handy und die Freizeitkultur des zwanzigsten Jahrhunderts fehlen bei Buck und sind in der IDS nachgetragen. Wer Buck heute benutzt, benutzt ihn als Gerüst, nicht als Kanon.

5. WordNet: dieselbe Idee, maschinenlesbar

Die zweite Quelle stammt aus der Computerlinguistik. **WordNet**, seit den achtziger Jahren in Princeton aufgebaut und von George Miller und Christiane Fellbaum geprägt, ordnet den englischen Wortschatz nicht alphabetisch, sondern nach Bedeutungsbeziehungen: Synonymie, Über- und Unterordnung, Teil-Ganzes.

Für uns interessant ist eine Nebenstruktur. Die Lexikografinnen und Lexikografen, die WordNet aufgebaut haben, arbeiteten in getrennten Dateien — den *lexicographer files*. Deren Namen sind faktisch eine grobe semantische Klassifikation und werden in der Computerlinguistik als **Supersenses** weiterverwendet: **26 Klassen für Nomen** (unter anderem `noun.animal`, `noun.artifact`, `noun.body`, `noun.food`, `noun.time`), dazu Klassen für Verben, Adjektive und Adverbien.

Der Wert dieser Liste liegt nicht darin, dass sie besser wäre als Buck. Er liegt darin, dass sie **von einer ganz anderen Seite her** gebaut wurde — von Leuten, die englische Bedeutungsrelationen für Maschinen beschreiben wollten, nicht Grundwortschatz für den Sprachvergleich. Wo zwei so verschiedene Systeme sich decken, ist die Kategorie vermutlich robust. Wo sie auseinandergehen, lohnt das Hinsehen.

Und genau da liegt ihre zweite Stärke: Buck ist stark bei den Dingen und schwach bei den Verben. WordNet hat für Verben eigene Klassen — Bewegung, Wahrnehmung, Kognition, Kommunikation, Besitz, soziales Handeln. Wer ein Vokabular mit fünfhundert Verben ordnen will, braucht das.

6. Der Farbttest: warum »bunt« keine Farbe ist

Es gibt ein Feld, an dem sich seit Jahrzehnten prüfen lässt, ob semantische Kategorien etwas Allgemeines haben oder bloß Konvention sind: die Farbwörter.

Brent Berlin und Paul Kay legten 1969 die These vor, dass Sprachen ihre **Grundfarbwörter** nicht beliebig wählen. Sie fanden im Englischen elf und behaupteten eine implikative Reihenfolge: Wer nur zwei Farbwörter hat, hat hell und dunkel; kommt ein drittes dazu, ist es rot; danach grün oder gelb, dann blau, dann braun — und erst zuletzt lila, rosa, orange, grau.

Die Universalitätsthese ist **stark umstritten** und in ihrer ursprünglichen Härte nicht zu halten; Anna Wierzbicka und Stephen Levinson haben ihr methodisch wie empirisch widersprochen, und spätere Erhebungen haben das Bild deutlich differenziert. Ich führe Berlin und Kay hier nicht als Beweis an, sondern wegen einer Unterscheidung, die davon unberührt bleibt und die man in der Praxis dauernd braucht:

Farbwörter und Farbeigenschaften sind nicht dasselbe. *Rot, blau, gelb, grün* benennen einen Farbton. *Hell, dunkel, bunt, durchsichtig* benennen keinen — sie beschreiben, wie ein Farbton auftritt. In Berlins und Kays Systematik ist das der Unterschied zwischen einem Grundfarbwort und einer Helligkeits- oder Sättigungsangabe.

Auf unserer Farbseite standen diese beiden Sorten gemischt, und niemandem fiel auf, warum die Seite unruhig wirkte. Der Grund war nicht das Layout. Der Grund war, dass zwei semantische Klassen dieselbe Fläche teilten. Seit sie getrennt stehen — die Farbtöne in einem Block, die Eigenschaften in einem zweiten —, liest sich die Seite von selbst.

Das ist die praktische Ausbeute einer theoretischen Unterscheidung, und sie ist typisch: Man braucht die Theorie nicht, um die Seite zu bauen. Man braucht sie, um zu *erklären*, warum die eine Fassung besser ist als die andere — und um es beim nächsten Mal gleich richtig zu machen.

7. Die zweite Achse: Wortart statt Bedeutung

Ein Vokabular ordnet nie nur nach Bedeutung. Es ordnet immer zugleich nach **Wortart**, und diese zweite Achse ist in der Unterstützten Kommunikation älter als jede Theorie über semantische Felder.

Sie geht auf Edith Fitzgerald zurück, eine gehörlose amerikanische Pädagogin, die in den zwanziger Jahren ein Schema entwickelte, mit dem gehörlose Kinder den Satzbau des Englischen sichtbar vor sich hatten: Wer? — Was tut? — Was? — Wo? — Wann? Der **Fitzgerald Key** war jahrzehntelang in einem Großteil der amerikanischen Gehörlosenschulen in Gebrauch und ist über Umwege zu dem geworden, was heute jede Symboltafel färbt: Nomen in einer Farbe, Verben in einer anderen, Eigenschaftswörter in einer dritten.

Wichtig ist, dass diese Achse **quer** zur semantischen liegt und nicht mit ihr verrechnet werden darf. *Essen* ist semantisch Nahrung und grammatisch je nach Schreibung Nomen oder Verb. *Rot* ist semantisch Farbe und grammatisch Adjektiv. Eine Seite kann nach der einen Achse vorbildlich sortiert sein und nach der anderen chaotisch.

In unseren Vokabularen sind deshalb beide Achsen gleichzeitig sichtbar: Die **Kategorie** entscheidet, auf welcher Seite ein Wort liegt; die **Wortart** entscheidet, in welcher Spalte und in welcher Farbe. Wer eine der beiden Achsen für die ganze Ordnung hält, baut zwangsläufig etwas Halbes.

8. Ein Prüfverfahren in drei Achsen

Damit ist das Werkzeug beisammen. Ein gewachsenes Vokabular lässt sich in drei Durchgängen prüfen, und jeder beantwortet eine andere Frage.

Achse A — Struktur. Bildet das System die anerkannten Bedeutungsfelder überhaupt ab? Dafür stellt man den eigenen Kategorien die Prototypen-Liste gegenüber und sucht nach Feldern ohne Entsprechung. Ein Feld ohne Liste ist eine Lücke; eine Liste ohne Feld ist entweder eine Eigenheit mit gutem Grund — oder ein Sammelbecken.

Hier liegt die erste Falle, und ich bin voll hineingelaufen. Mein erster Durchgang ordnete auf der Ebene der *Hauptgruppen* zu und meldete drei unbelegte Felder: Handeln und Technik, Konflikt und Schutz, Technik und Medien. Alle drei waren belegt — durch »Arbeit – Werkzeug«, »Welt – Frieden« und »Stadt – Internet & Apps«. Die Lücke war meine eigene Grobheit. Merkregel: **erst auf Listenebene zuordnen, dann Lücken behaupten.**

Achse B — Wörter. Liegt das einzelne Wort in der Kategorie, die zu ihm passt? Das lässt sich nicht durch Nachdenken beantworten, denn man würde wieder die eigene Zuordnung mit dem eigenen Kopf prüfen. Man braucht ein **unabhängiges Signal**.

Bei Symbolvokabularen liegt eines bereit, das oft übersehen wird: die **Ordnersystematik der Symbolsammlung**. METACOM ordnet seine Symbole in rund sechzig Sachgebiete — Tiere, Körperteile, Lebensmittel, Werkzeug, Kleidung. Diese Einteilung stammt von der Autorin der Symbole und nicht von uns; sie ist damit genau das, was die Prüfung braucht. Weicht das Sachgebiet des Symbols vom Feld der Liste ab, ist das ein Grund hinzusehen.

Achse C — Homogenität. Ist eine Liste überhaupt ein semantisches Feld — oder eine Situation? Das misst man an der Streuung: Wie groß ist der Anteil des häufigsten Feldes innerhalb der Liste? Liegt er hoch, ist die Liste ein Feld. Liegt er niedrig, ist sie nach einem **Ereignis** gebaut.

Diese dritte Achse ist mir erst beim Auswerten aufgefallen, und sie ist die interessanteste. Denn eine niedrige Homogenität ist **kein Fehler**. Eine Liste »Sich vorstellen« enthält Sprachen, Familienstand und Adresse — drei völlig verschiedene Felder, aber eine einzige Situation. Genau so bauen erfahrene Systeme ihre Aktivitätsseiten, und genau dafür gibt es

gute Gründe: Kinder gruppieren eher nach Ereignis als nach Oberbegriff (siehe Schwestertext). Achse C sagt nicht, was falsch ist. Sie sagt, **welche Listen nach welchem Prinzip gebaut sind** — und das sollte man wissen, bevor man an ihnen etwas ändert.

9. Warum eine rohe Fehlerquote lügt

Jetzt kommt der Teil, den ich für den wichtigsten dieses Textes halte.

Der erste vollständige Durchgang über unseren Bestand meldete **27,8 % Abweichung**. Wer diese Zahl in einen Bericht schreibt, hat einen Skandal. Wer sie sich ansieht, hat ein Messproblem.

Denn fast alles davon war Eigenrauschen, in zwei Sorten.

Erstens: benachbarte Felder, die sich sachlich überschneiden. Eine Zahnbürste ist Körperpflege *und* Gesundheit. Ein Datum ist Zeit *und* Zahl. Ein Land ist Geografie *und* Politik. Keine dieser Kreuzungen ist ein Fehler; sie zeigt nur, dass Felder keine disjunkten Behälter sind — was Trier schon wusste. Man muss solche Paare vorher benennen und als verträglich behandeln. Das allein brachte die Quote von 27,8 % auf **11,4 %**.

Zweitens: Sachgebiete ohne Domänensignal. Manche Symbolordner sind querliegend. Ein *Beruf* kommt in jedem Feld vor — der Zahnarzt in der Gesundheit, der Lehrer in der Schule, der Koch beim Essen. Ein Ordner »Stadt und Verkehr« mischt Gebäude und Wege, also kollidiert er mit jedem Lokal und jedem Amt. Ein *Buch* handelt von irgendetwas. Solche Ordner tragen keine Information über das Feld und müssen aus der Prüfung heraus. Gemessen trieben fünf davon 69 der verbliebenen 199 Abweichungen. Danach: **8,2 %**.

Drei Zahlen für denselben Bestand — 27,8 %, 11,4 %, 8,2 % —, und nur die letzte sagt etwas über das Vokabular. Die beiden ersten sagen etwas über den Prüfer.

Daraus folgt eine Regel, die ich für übertragbar halte: **Eine automatische Prüfung ist erst dann fertig, wenn ihre Fehlalarme benannt und begründet ausgeschlossen sind.** Solange das nicht geschehen ist, ist die Quote keine Aussage über den Gegenstand. Und der Ausschluss muss *inhaltlich* begründet sein — nicht dadurch, dass man so lange filtert, bis die Zahl gefällt.

Das Ergebnis solcher Prüfungen ist im Übrigen nie eine Fehlerliste, sondern eine **Durchsichtsliste**. Ein Wort darf mit Absicht quer liegen: Der Fisch im Tierlexikon und der Fisch auf der Speisekarte sind zwei verschiedene Dinge, und beide gehören ins Vokabular. Wer solche Fälle »korrigiert«, verschlechtert das System.

10. Was dabei herauskam

Für unseren eigenen Bestand — 31 Hauptgruppen, 241 Listen, rund 8.600 Einträge — lautet das Ergebnis:

- **Alle 25 Prototyp-Felder sind belegt.** Es fehlt kein Bedeutungsbereich.
- **8,2 % Abweichung** auf Wortebene nach begründeter Filterung, und der weitaus größte Teil davon ist sachlich richtig: Ein Schulfach »Sport« trägt zu Recht ein Sportsymbol, Arbeitskleidung gehört in die Fabrikliste.
- **20 von 128 auswertbaren Listen sind semantisch heterogen** — also nach einer Situation gebaut. Das ist gewollt.

Was übrig blieb, war eine knappe Handvoll echter Funde. Drei Zeitadverbien standen im Sammeltopf, obwohl es für sie eine eigene Liste gab. In der Küchenliste steckten ein Briefumschlag und eine Spielkiste, beide Dubletten von Wörtern, die anderswo längst zu Hause waren. Und ein *Korb* trug das Symbol eines **Basketballkorbs** — die Symbolsammlung führt beides unter demselben Suchwort.

Diese Ausbeute wirkt mager, und das ist die eigentliche Nachricht: **Der Bestand ist gut kategorisiert.** Ein Audit, das nichts findet, ist kein vergeudetes Audit — es ist der Beleg, den man vorher nicht hatte. Und es hinterlässt ein Werkzeug, das beim nächsten Wachstumsschub wieder läuft.

11. Wie wir es bei AssistUK halten

Aus alldem sind fünf Arbeitsregeln geworden.

Wir ordnen nach zwei Achsen und halten sie auseinander. Die Bedeutung entscheidet über die Seite, die Wortart über Spalte und Farbe. Wo beides zusammenfällt, ist es Zufall, kein Prinzip.

Wir erlauben beides: Felder und Situationen. Ein Vokabular, das nur Oberbegriffe kennt, passt zu keinem Alltag; eines, das nur Situationen kennt, lässt sich nicht durchsuchen. Neu ist, dass wir jetzt *wissen*, welche unserer Listen zu welcher Sorte gehört, statt es zu ahnen.

Wir messen gegen ein System von außen. Die Prototypen-Liste ist nicht unsere; das ist ihr ganzer Wert. Ebenso die Sachgebiete der Symbolsammlung.

Wir behandeln jede automatische Prüfung als Durchsichtsliste. Keine Zuordnung wird ohne menschliche Entscheidung geändert, und ein bewusst quer liegendes Wort bleibt liegen.

Wir schreiben die Begründung dazu. Warum eine Liste geteilt wurde, warum ein Wort umgezogen ist, warum ein Feldpaar als verträglich gilt — das steht in den Werkzeugen selbst. Ein Kategoriensystem, dessen Gründe nur im Kopf einer Person liegen, ist beim nächsten Umbau verloren.

12. Grenzen dieses Textes

Ich beschreibe ein **Verfahren**, keine Wirksamkeitsstudie. Dass ein nach diesen Feldern geordnetes Vokabular besser bedienbar wäre als ein anders geordnetes, behaupte ich nicht; dazu gibt es die Befunde im Schwestertext, und die sind differenzierter, als es einem lieb ist.

Die Prototypen-Liste in Abschnitt 8 ist eine **Zusammenstellung**, keine Quelle: Sie verbindet Bucks Kapitel mit den WordNet-Klassen und ergänzt vier Felder für die moderne Lebenswelt und die Funktionswörter. Diese Zusammenstellung ist meine Verantwortung, und sie ist diskutabel — ein Kapitel »Krieg und Jagd« habe ich zu »Konflikt und Schutz« umgewidmet, weil ein Kommunikationsvokabular andere Aufgaben hat als eine indogermanistische Wortstudie.

Der unabhängige Maßstab in Achse B ist **eine bestimmte Symbolsammlung**. Wer mit anderen Symbolen arbeitet, braucht ein anderes Signal — oder muss auf Achse B verzichten und sich auf A und C beschränken.

Zu Berlin und Kay: Ich benutze ihre **Unterscheidung** zwischen Grundfarbwort und Farbeigenschaft, nicht ihre Universalitätsthese. Letztere ist umstritten, und ich stütze auf sie kein Argument.

Zuletzt: Alle Zahlen in diesem Text stammen aus **einem** Bestand, meinem eigenen. Die Größenordnungen sind nicht übertragbar; das Verfahren, so hoffe ich, schon.

13. Fazit in fünf Sätzen

Ein gewachsenes Kategoriensystem hat niemand entworfen, und es entscheidet trotzdem täglich darüber, ob jemand sein Wort findet. Prüfen lässt es sich nur gegen einen Maßstab, der nicht aus dem eigenen Kopf stammt — für Bedeutungsfelder liefern Buck und die IDS einen, für Wortarten der Fitzgerald Key, für die Gegenprobe am einzelnen Wort die Sachgebiete der Symbolsammlung. Wer so prüft, muss seine eigenen Fehlalarme zuerst benennen, sonst misst die Quote den Prüfer statt den Bestand. Nicht jede Liste ist ein Bedeutungsfeld: Manche sind nach einer Situation gebaut, und das ist richtig so — man sollte nur wissen, welche. Und ein Audit, das am Ende wenig findet, hat seinen Zweck erfüllt, denn es liefert den Beleg, den man vorher nicht hatte.

14. Literatur (mit Prüfvermerk)

Alle Quellen am 22.08.2026 auf Existenz und Kernaussage geprüft (✓); ○ kennzeichnet Angaben, die ich nicht am Original nachgeschlagen habe.

1. ✓ **Intercontinental Dictionary Series (IDS)**. Hrsg. Mary Ritchie Key (Begründerin) & Bernard Comrie. Leipzig: Max-Planck-Institut für evolutionäre Anthropologie. — Wortliste mit **1.310 Bedeutungen in 22 Kapiteln**, nach dem thematischen Aufriss und der Nummerierung von Buck (1949) gebaut und von dessen rund 1.200 Einträgen erweitert; fehlende Wörter bleiben leer. (Am 22.08.2026 auf ids.clld.org geprüft.)

2. ✓ **Buck, C. D. (1949).** *A Dictionary of Selected Synonyms in the Principal Indo-European Languages*. Chicago: University of Chicago Press. — Die Ursprungsgliederung in 22 semantische Kapitel. (Existenz und Rolle über die IDS bestätigt; das Werk selbst liegt mir nicht vor.)
3. ✓ **Trier, J. (1931).** *Der deutsche Wortschatz im Sinnbezirk des Verstandes. Die Geschichte eines sprachlichen Feldes*. Heidelberg: Winter. — Wortfeldtheorie: Die Bedeutung eines Wortes entsteht durch die Abgrenzung gegen seine Feldnachbarn. ○ (Standardreferenz; Seitenangaben nicht geprüft.)
4. ✓ **Lyons, J. (1977).** *Semantics*. 2 Bände. Cambridge: Cambridge University Press. — Lexikalische Felder und Sinnrelationen als beschreibbares Instrumentarium. ○ (Standardreferenz; Seitenangaben nicht geprüft.)
5. ✓ **Fellbaum, C. (Hrsg.) (1998).** *WordNet: An Electronic Lexical Database*. Cambridge, MA: MIT Press. — Aufbau und Relationen von WordNet; Grundlage der Supersense-Klassifikation. Vgl. auch Miller, G. A. (1995), *WordNet: A Lexical Database for English*, *Communications of the ACM* 38(11).
6. ✓ **WordNet — lexicographer files / Supersenses.** Princeton University. — **26 Klassen für Nomen**, dazu eigene Klassen für Verben, Adjektive und Adverbien; in der Computerlinguistik als grobe semantische Klassifikation weiterverwendet. (Die Zahl 26 für die Nomenklassen am 22.08.2026 bestätigt; die Verbklassen ○ nach der Standard-`lexnames`-Liste, nicht am Original nachgezählt.)
7. ✓ **Berlin, B. & Kay, P. (1969).** *Basic Color Terms: Their Universality and Evolution*. Berkeley: University of California Press. — Grundfarbwörter und ihre behauptete implikative Reihenfolge. **Die Universalitätsthese ist stark umstritten** (u. a. Wierzbicka, Levinson) und wird hier nicht als Beleg verwendet; übernommen ist allein die Unterscheidung zwischen Grundfarbwort und Farbeigenschaft.
8. ✓ **Fitzgerald, E. (1926).** *Straight Language for the Deaf: A System of Instruction for Deaf Children*. — Ursprung des »Fitzgerald Key«, der Satzbau über feste Positionsspalten sichtbar macht und über die Gehörlosenpädagogik zur Wortart-Farbkodierung in der Unterstützten Kommunikation geführt hat. ○ (Erscheinungsjahr in der Literatur uneinheitlich mit 1926 oder 1929 angegeben; Edith Mansford Fitzgerald 1877–1940.)
9. ✓ **Wittgenstein, L. (1953).** *Philosophische Untersuchungen*. — Familienähnlichkeit: Die Dinge, die wir »Spiel« nennen, teilen kein einziges gemeinsames Merkmal. Gedanklicher Vorläufer der Prototypentheorie. ○ (Standardreferenz; Paragrafenangaben nicht geprüft.)
10. ✓ **Fallon, K. A., Light, J. C. & Achenbach, A. (2003).** The semantic organization patterns of young children: Implications for augmentative and alternative communication. *Augmentative and Alternative Communication*, 19(2), 74–85. DOI: 10.1080/0743461031000112061. — Ob die von Erwachsenen gebauten Anordnungen der kognitiven Organisation junger Kinder entsprechen, war offen; die Studie prüft das. (Zitat und Fundstelle am 22.08.2026 bestätigt; Einzelergebnisse referiert der Schwestertext.)
11. ✓ **Tiedemann, L. (2026).** Wo die Wörter wohnen — Kategorien, Prototypen und Aktivitätsseiten. AssistUK-Grundlagentext. — Schwestertext: Prototypentheorie, Basisebene, Ereignisschemata, Kategorienkatalog, Aktivitätsseiten. (Hausintern; Quellenapparat dort, darunter Rosch & Mervis 1975 und Rosch et al. 1976.)
12. ✓ **Tiedemann, L. (2026).** Die zweihundert Wörter, die fast alles sagen. AssistUK-Grundlagentext. — Kern- und Randvokabular, Positions Konstanz, Farbkonvention. (Hausintern; Quellenapparat dort.)
13. ✓ **Tiedemann, L. (2026).** Die Anordnung schlägt die Farbe. AssistUK-Grundlagentext. — Wortart-Achse, Rastergröße, Farbe. (Hausintern; Quellenapparat dort.)

Die fachliche Position, Auswahl und Verantwortung liegen bei mir. — L. T.