

既知0文字からの解読手順

本文出典：フロストヘイヴン / Step 1 - 10

Step 1：「ここに」の反復から K・O・N・I を導く

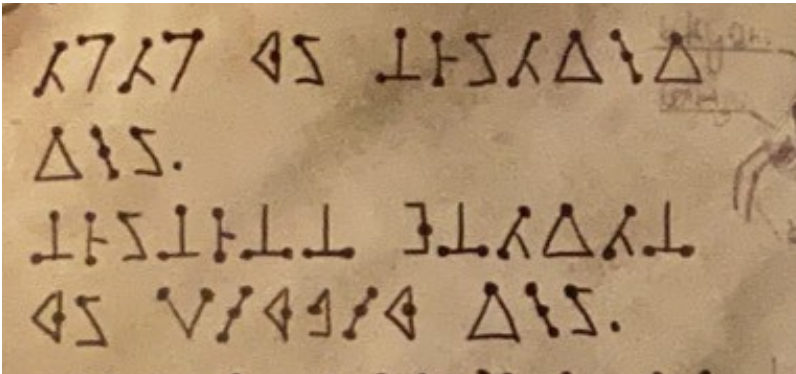
開始時点：判明 0 / 未判明 26

A B C D E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：ガイドの日本語訳と字形列の照合

使う手がかりは本文画像と、ガイドにある冒頭の訳「ここに力あり。地中深くに源泉あり」。読みの対応は0文字から始める。小さなIDは字形を識別するためのラベルで、読みを教える情報としては使わない。解釈は古文書として自然な文体を優先し、現代口語との違いだけで直さない。現代口語でしか読めない箇所は、字形・文法・文脈を再確認し注記する。



本文画像・左ページ冒頭（最初の二節）

まず、日本語の音をアルファベットへ置き換えたローマ字ではないかと仮定する。「ここに」を KOKO NI とすると、冒頭6字の反復が一致する。単語の区切りは訳との照合から補う。

字形 → このStepの開始時点の読み（小数字＝固定ID）

--	--	--	--	--	--

→ KOKO NI / ここに

1字目と3字目は同じID 4、2字目と4字目は同じID 21。K O K O の繰り返しと合うため、4=K、21=Oを置き、続く3=N、23=Iを導く。この段階では読み方の仮説であり、次の訳にも同じ対応を当てて確かめる。

このStepでの採用：03=N 04=K 21=O 23=I / 終了時点：4文字

Step 2 : 「力あり」の文字数と反復から C・H・A・R

開始時点：判明 4 / 未判明 22

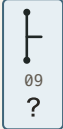
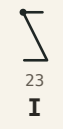
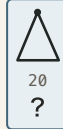
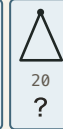
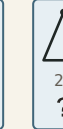
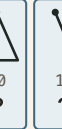
A B C D E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：ガイドの日本語訳と字形列の照合

ガイドの続き「力あり」を、本文の次の10字に合わせる。前のStepで得たIとKが、候補 CHIKARA ARI の同じ位置に現れる。

字形 → このStepの開始時点の読み（小数字＝固定ID）

									
?	?	I	K	?	?	?	?	?	I

→ CHIKARA ARI / 力あり

同じID 20が現れる3か所はすべてA、ID 12の2か所はすべてRになる。先頭2字をC・Hとすれば、前のStepのI・Kとも矛盾しない。7=C、9=H、20=A、12=Rを加える。

「ち」をTIと書く TIKARA ARI では9字になり、本文の10字と合わない。CHIなら10字になる。この比較は、字形1個をアルファベット1文字に置換する仮定の下で行う。訳の意味だけで読みを決めず、文字数と反復も合わせる。

ここまでに8文字。上のA～Z一覧では、取り消し線の文字はすでに別のIDに対応している。字形とアルファベットが一对一なら、それらを新しいIDへ重ねて割り当てることはできない。

このStepでの採用：07=C 09=H 12=R 20=A / 終了時点：8文字

Step 3 : 二節目の訳から U・F・G・E・S を導く

開始時点 : 判明 8 / 未判明 18

A B C D E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線=判明済み(候補から除外) 太字=未判明 枠=このStepで導く文字

根拠 : ガイドの日本語訳と字形列の照合

次の訳「地中深くに源泉あり」を CHICHUU FUKAKU NI GENSEN ARI として照合する。長音をUUと書くと、本文で同じID 6が2回続く箇所合う。

字形 → このStepの開始時点の読み (小数字=固定ID)

C	H	I	C	H	?	?	?	?	K	A	K	?

→ CHICHUU FUKAKU / 地中深く

先頭のCHIはStep 2で得た3文字と同じ並び。6=Uなら CHICHUU と FUKAKU の両方で一致する。「ふ」をHUとする案では、別のID 9に対応するHをID 19にも割り当ててしまう。そこで19=Fを採用する。

字形 → このStepの開始時点の読み (小数字=固定ID)

N	I	?	?	N	?	?	N	A	R	I

→ NI GENSEN ARI / に源泉あり

GENSENの3字目・6字目は既知のN、2字目・5字目は同じ未知ID 13。11=G、13=E、15=Sを置けば、前後のNI・ARIも含めて訳に一致する。これで13文字を、最初の表の読みを使わずに導けた。

以後はガイドの訳がない。残る候補は B・D・J・L・M・P・Q・T・V・W・X・Y・Z。ここからは単語と文脈による推定を、別の出現箇所確かめながら進める。

このStepでの採用 : 06=U 11=G 13=E 15=S 19=F / 終了時点 : 13文字

Step 4 : 「我らは」「を」 から W を導く

開始時点：判明 13 / 未判明 13

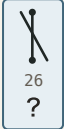
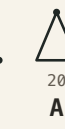
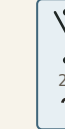
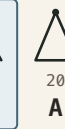
A B C D E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：本文の複数の語と文脈

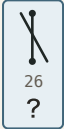
左ページの続きにある同じ未知ID 26を、まず二つの語で読む。既知13文字だけを入れると「? ARERA ? A」になる。

字形 → このStepの開始時点の読み（小数字＝固定ID）

							
26	20	12	13	12	20	26	20
?	A	R	E	R	A	?	A

→ WARERA WA / 我らは

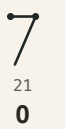
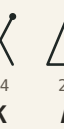
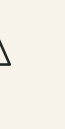
字形 → このStepの開始時点の読み（小数字＝固定ID）

	
26	21
?	O

→ WO / を（同じ左ページの後続箇所）

26＝Wなら、WARERA WAとWOの両方が日本語として読める。助詞「は」「を」はここではWA・WOと書かれているとみなす。Wを加え、残る未判明は12文字になる。

字形 → このStepの開始時点の読み（小数字＝固定ID）

				
18	21	04	04	20
?	O	K	K	A

→ ?OKKA / 先頭のID 18は保留

この直前の「?OKKA」は、まだ意味だけでMと決めない。KOKKA（国家）という案は、KがID 4に対応済みでID 18とは別字形なので除外できる。一方、Tはまだ未判明なので、この時点で候補から外してはいけない。次の語も合わせて判断する。

このStepでの採用：26＝W / 終了時点：14文字

Step 5 : 「目下」「三つ」「建てた」を合わせて M・T

開始時点：判明 14 / 未判明 12

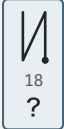
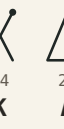
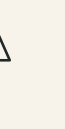
A B C D E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：本文の複数の語と文脈

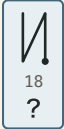
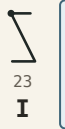
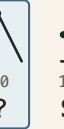
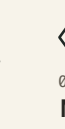
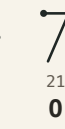
保留したID 18は、次の「?I??SU」の先頭にも現れる。その語では別の未知ID 10が2回連続する。二つの語を同時に読める対応を探す。

字形 → このStepの開始時点の読み（小数字＝固定ID）

				
18	21	04	04	20
?	O	K	K	A

→ MOKKA / 目下

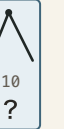
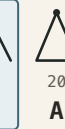
字形 → このStepの開始時点の読み（小数字＝固定ID）

							
18	23	10	10	15	06	03	21
?	I	?	?	S	U	N	O

→ MITTSU NO / 三つの

18=M、10=Tとすれば MOKKA と MITTSU がともに成立し、連続するID 10がTTに対応する。「目下、我らは三つの…」とつながる。ID 18をTにする案では、別のID 10もTと読むことはできず、MITTSUという読みと両立しない。

字形 → このStepの開始時点の読み（小数字＝固定ID）

					
10	20	10	13	10	20
?	A	?	E	?	A

→ TATETA / 建てた

同じID 10が3回現れる語もTATETAと読める。この段階の未使用12文字を、ID 18・10へ重複なく割り当てる132通りをすべて列挙した。ID 10の12候補を比べ、Tを採用した後のID 18の11候補も比較すると、3語が文脈に合うM・Tを選べる。全候補の比較は巻末の「補足：M・Tの総当たり」を参照。

このStepでの採用：10=T 18=M / 終了時点：16文字

Step 6 : 「研究施設」「奴ら」から Y を導く

開始時点：判明 16 / 未判明 10

A B C D E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：本文の複数の語と文脈

左ページの長い語では、ここまでの16文字を入れると未知の字形はID 16だけになる。

字形 → このStepの開始時点の読み（小数字＝固定ID）

														
K	E	N	K	?	U	U	S	H	I	S	E	T	S	U

→ KENKYUUSHISETSU / 研究施設

字形 → このStepの開始時点の読み（小数字＝固定ID）

						
?	A	T	S	U	R	A

→ YATSURA / 奴ら

16=Yなら KENKYUUSHISETSU と YATSURA の両方が読める。前のStepとつなげると「三つの研究施設を建てた」となる。右ページのHAYASUGIRU（早すぎる）、HITSUYOU（必要）でも同じ対応を照合する。

このStepでの採用：16=Y / 終了時点：17文字

Step 7 : 「同僚ら」「通り」「のだ」から D を導く

開始時点：判明 17 / 未判明 9

A B € **D** E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：本文の複数の語と文脈

左ページのTATETA.の後と、別の短い語に、同じ未知ID 22が現れる。

字形 → このStepの開始時点の読み（小数字＝固定ID）

22 ?	21 O	06 U	12 R	16 Y	21 O	06 U	12 R	20 A

→ DOURYOURA / 同僚ら

字形 → このStepの開始時点の読み（小数字＝固定ID）

22 ?	21 O	06 U	12 R	23 I

→ DOURI / 通り

字形 → このStepの開始時点の読み（小数字＝固定ID）

03 N	21 O	22 ?	20 A

→ NODA / のだ

22=Dという同じ対応で3例が成立する。DOURYOURAは画像の字形に沿ったつづりを保持する。意味の似た別の語に置き換えない。

このStepでの採用：22=D / 終了時点：18文字

Step 8 : 「すべて」「展望」「二倍」から B を導く

開始時点：判明 18 / 未判明 8

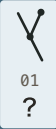
A B € Đ Ē F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：本文の複数の語と文脈

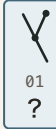
右ページの未知ID 1を、三つの語で照合する。このStepの開始時点では18文字が判明済みで、未判明の候補は8文字に減っている。

字形 → このStepの開始時点の読み（小数字＝固定ID）

					
15	06	01	13	10	13
S	U	?	E	T	E

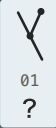
→ SUBETE / すべて

字形 → このStepの開始時点の読み（小数字＝固定ID）

					
10	13	03	01	21	06
T	E	N	?	O	U

→ TENBOU / 展望

字形 → このStepの開始時点の読み（小数字＝固定ID）

				
03	23	01	20	23
N	I	?	A	I

→ NIBAI / 二倍

1=Bを当てはめると3語とも成立する。左ページのNARABAも同じBで読める。原文はNARABAとして転記し、現代口語に合わせてNAREBAへ書き換えない。

このStepでの採用：01=B / 終了時点：19文字

Step 9 : 「ゼロとなる」 から Z を導く

開始時点 : 判明 19 / 未判明 7

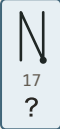
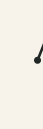
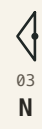
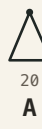
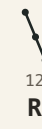
A B C D E F G H I J K L M N O P Q R S T U V W X Y Z

取り消し線=判明済み（候補から除外） 太字=未判明 枠=このStepで導く文字

根拠：本文の1例による判断

左ページの残る語は「?ERO TO NARU」。上の一覧で残る候補はJ・L・P・Q・V・X・Z。HEROのHはすでにID 9に対応済みなので、この別字形には使えない。

字形 → このStepの開始時点の読み（小数字=固定ID）

									
17 ?	13 E	12 R	21 O	10 T	21 O	03 N	20 A	12 R	06 U

→ ZERO TO NARU / ゼロとなる

17=ZならZERO TO NARUと読める。ただし、この本文ではZEROの1例による判断であり、複数の語で確認できた文字より根拠が弱い。対応にはこの制約を記録し、新しい文章で再確認する。

これで本文に現れる20種類の字形を読める。残りのJ・L・P・Q・V・Xについては、まだ個別の字形を割り当てる根拠がない。

このStepでの採用：17=Z / 終了時点：20文字

Step 10：全体を読み直し、20文字と未判明6文字を保存

開始時点：判明 20 / 未判明 6



取り消し線＝判明済み（候補から除外） 太字＝未判明 枠＝このStepで導く文字

根拠：ここまで導いた対応の再確認

0文字から開始し、ガイドの訳から13文字、本文の文脈から7文字を導いた。各Stepで同じIDに同じアルファベットを使い、一対一対応に矛盾しないことを確認する。

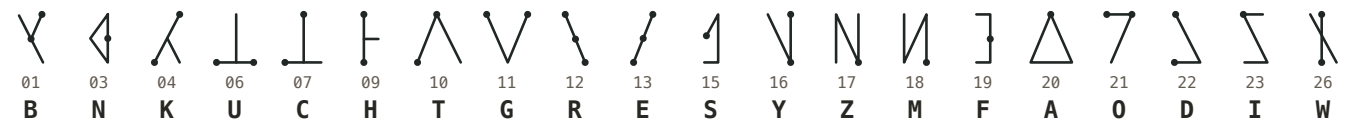
字形 → このStepの開始時点の読み（小数字＝固定ID）

26	20	10	20	15	09	23	03	21	10	13	03	01	21	06
W	A	T	A	S	H	I	N	O	T	E	N	B	O	U

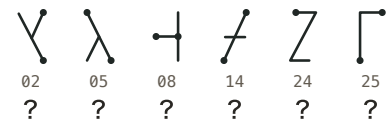
→ **WATASHI NO TENBOU / 私の展望（左末尾と右冒頭がつながる）**

末尾のWARERA NIWA NIBAI NO SHISETSU GA HITSUYOU NANODA.は「我らには二倍の施設が必要なのだ」と読める。先に読んだ「三つの研究施設」という内容ともつながる。

この版で導いた20文字



未判明の6字形：候補は J・L・P・Q・V・X



ID 2・5・8・14・24・25の個別対応は未判明のまま残す。これらの字形は照合用一覧にあるが、今回の本文では登場を確認できない。日本語訳とローマ字の照合、一対一対応という前提の下で再構成した手順であり、ガイドなしの暗号文だけから一意に解けると示したものではない。

このStepでの採用：新しい対応の追加なし / 終了時点：20文字

解読全文：アルファベット（ローマ字）

字形から読んだつづりを保存する。原画像の改行は文ごとにまとめ、空白を整理した。左ページ末尾のWATASHIと右ページ冒頭のNO TENBOUは続けて読む。[1]～[3]は注意箇所、本文の文字ではない。

KOKO NI CHIKARA ARI. CHICHUU FUKAKU NI GENSEN ARI.

MOKKA WARERA WA MITTSU NO KENKYUUSHISETSU WO TATETA. DOURYOURA NO OMOWAKU
DOURI[2] NI NARABA[3] ZERO[1] TO NARU.

YATSURA NIWA WATASHI NO TENBOU GA MIENU NODA. SUBETE GA KAWATTE SHIMAU KANOUSEI
MO ARU GA, OSORERU NIWA HAYASUGIRU.

WARERA NIWA NIBAI NO SHISETSU GA HITSUYOU NANODA.

読みの根拠・表記上の注意

- [1] ZEROのZ（ID 17）は、この本文では1例だけからの推定。語としては「ゼロ」と読めるが、複数語で照合できた文字より根拠が弱く、新しい資料で再確認する。
- [2] 字形の転記はDOURI。日本語として「どおり」と整える場合も、原文をDOORIに書き換えない。
- [3] 字形の転記はNARABA。文語の「なる」の未然形「なら」＋仮定の「ば」と解釈し、漢字を含む文章でも「にならば」を保持する。NAREBAに書き換えない。

解読全文：ひらがな

原文の音とつづりを追えるようにした読み。ここではDOURIを「どうり」、NARABAを「ならば」のまま残す。WAは「わ」、WOは通常のかな対応で「を」と書く。分かち書き・句読点は読みやすさのための整理。

こ こ に ち か ら あ り 。 ち ち ゅ う ふ か く に げ ん せ ん あ り 。

も っ か わ れ ら わ み っ つ の け ん き ゅ う し せ つ を た て た 。 ど う り ょ う ら の お も わ く ど う り **[2]** に な ら ば **[3]** ぜ ろ **[1]** と な る 。

や つ ら に わ わ た し の て ん ぼ う が み え ぬ の だ 。 す べ て が か わ っ て し ま う か の う せ い も あ る が 、 お そ れ る に わ は や す ぎ る 。

わ れ ら に わ に ば い の し せ つ が ひ つ よ う な の だ 。

読みの根拠・表記上の注意

- **[1]** 「ぜろ」はZの推定を引き継いだ読み。ひらがなにしても根拠の強さは変わらない。
- **[2]** 「どうり」はDOURIを追った表記。次ページでは、文脈上の「思惑どおり」に合わせて「どおり」とする。
- **[3]** 「に ならば」はNI NARABAのまま。文語の仮定表現として解釈できるため、次ページでも保持する。現代語での参考の言い換えは「になれば」であり、原文の訂正ではない。
- 助詞も音を追って「われら わ」「やつら にわ」とした。次ページの文章では標準的な表記の「我らは」「奴らには」へ整える。

解読全文：漢字を含む文章（古風な文体を保持）

古文書として自然な文体を保持し、漢字・助詞表記・句読点を整えた読み。[1]は根拠が弱い字の推定、[2]はかな遣いの整理、[3]は原文を保持する文法上の解釈を示す。

ここに力あり。地中深くに源泉あり。

目下、我らは三つの研究施設を建てた。同僚らの思惑どおり[2]になれば[3]、ゼロ[1]となる。

奴らには私の展望が見えぬのだ。すべてが変わってしまう可能性もあるが、恐れるには早すぎる。

我らには二倍の施設が必要なのだ。

読みの根拠・表記上の注意

- [1] Z → 「ゼロ」 → 「ゼロ」：Z=ID 17はZEROの1例による推定。何がゼロになるかは本文で明示されていないため、主語や対象を補っていない。
- [2] DOURI → 「どうり」 → 「どおり」：文脈から「通り」の意味と解釈し、かな遣いを整えた。ローマ字からの機械的な変換結果ではない。
- [3] NI NARABA → 「に ならば」 → 「になれば」：「なる」の未然形「なら」＋「ば」による文語の仮定表現と解釈する。現代語の参考訳は「になれば」。前版の「文法的に不自然なので補正する」という判断は撤回し、原文を保持する。
- WA・NIWA → 「わ・にわ」 → 「は・には」：助詞の標準表記へ変更。WOは「を」。長音は「けんきゅう」「どうりょう」「てんぼう」などの通常のかな表記にした。
- 漢字は文脈から選んだ表記（例：DOURYOURA＝同僚ら、TENBOU＝展望、TATETA＝建てた）。字形そのものが漢字を指定するわけではない。文の意味を詳しく補う意訳は加えていない。
- 「建てた」「必要なだ」など口語的な形も原文にある。古風だけを理由に書き換えず、字形と文脈を再確認して保持する。文章全体が特定の時代の純粋な文語であると断定するものではない。

文法資料：デジタル大辞泉「成る」 / 「ば」（コトバンク）

補足：M・Tの総当たり（Step 5）

Step 5開始時点の未使用文字はB・D・J・L・M・P・Q・T・V・X・Y・Z。ID 18とID 10は別字形なので、重複を除く12×11=132通りを列挙する。

3語は「18 O K K A」「18 I 10 10 S U」「10 A 10 E 10 A」。まずID 10だけで決まる3語目を比較し、次にID 18を比べると、132通りを次の2表で確認できる。

① ID 10：12候補すべて

ID 10	10 A 10 E 10 A
B	BABEBA
D	DADEDA
J	JAJEJA
L	LALELA
M	MAMEMA
P	PAPEPA
Q	QAQEQA
T	TATETA
V	VAVEVA
X	XAXEXA
Y	YAYEYA
Z	ZAZEZA

② ID 10=Tとして残る11候補

ID 18	18 O K K A	18 I T T S U
B	BOKKA	BITTSU
D	DOKKA	DITTSU
J	JOKKA	JITTSU
L	LOKKA	LITTSU
M	MOKKA	MITTSU
P	POKKA	PITTSU
Q	QOKKA	QITTSU
V	VOKKA	VITTSU
X	XOKKA	XITTSU
Y	YOKKA	YITTSU
Z	ZOKKA	ZITTSU

候補を外す理由

- ① 「…を建てた」と読めるTATETAを採用する。他の11候補は、この文脈で動作を表す自然な語として読めない。ID 18の選び方を変えても3語目は変わらないため、ここで121通りを外せる。
- ② TをID 10へ割り当てたので、ID 18には残る11文字を試す。MならMOKKA（目下）とMITTSU（三つ）がともに成立する。DOKKAやYOKKAは1語だけなら読みを考えられるが、DITTSU・YITTSUまで同時に文脈へ合わせられない。残る候補も2語の組として読めず、Mを採用する。

結果：18=M、10=T。「目下、我らは三つの…を建てた」とつながる。

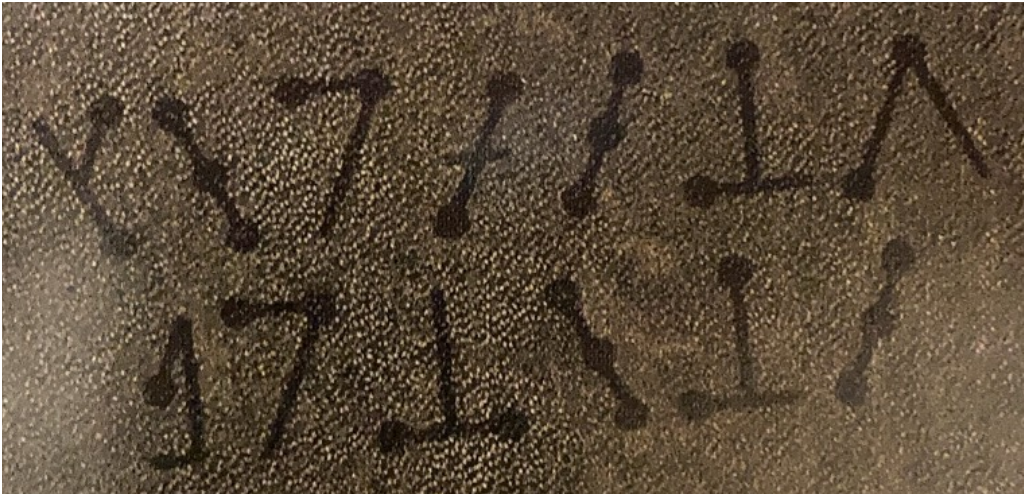
全132通りの文字列はstep5-candidates.csvに保存。候補の列挙と一対一条件は自動検査し、語・文脈の自然さは古風な文体も考慮して比較した。

続き：本の表紙を読む

本文で導いた20文字を表紙へ適用する。プレイヤーから「PROJECT SOURCEではないか」という仮説が提示されたため、字形と別候補の両方を検証した。

開始時点：採用20文字・未判明6文字 / 枠=今回の暫定採用P

A B C D E F G H I J K L M N O P Q R S T U V W X Y Z



cover_page.jpg (原画像の文字部分、補筆なし)

字形 → 表紙を調べる前の読み (小数字=固定ID)

 02 ?	 12 R	 21 O	 14 ?	 13 E	 07 C	 10 T
 15 S	 21 O	 06 U	 12 R	 07 C	 13 E	

→ ?RO?ECT / SOURCE

下段はすべて既知文字でSOURCE。上段のR・O・E・Cは下段にも現れ、点・線の向きと反復を照合できる。本文がローマ字でも、表題は英語として読む。

ID 02は右上端と右下端に点がある分岐形 (B=ID 01は中央に点)。ID 14は左下がりの斜線と中央横棒、両端の点として転記した。横棒は薄く、隣のE=ID 13の中央点とは区別するが、鮮明な再出現でも確認する。

有力な読み：PROJECT SOURCE

読み：ぷろじえくと そーす / 参考訳：「源泉計画」。SOURCEが固有の計画名か、何の源を指すかは表紙だけでは決められない。

別候補PROVECT SOURCEが残る。ID 02=Pは暫定採用、ID 14=Jは有力候補とし、Vを保留する。表紙だけによる一意の確定とはしない。

表紙の別候補：30通りを照合する

未使用文字J・L・P・Q・V・Xを、別字形ID 02・14へ重複なく割り当てる。6×5=30通り。下段はすべてSOURCEのまま。表の※は照合した単語一覧への完全一致を示す。

JR0LECT	JR0PECT	JR0QECT	JR0VECT	JR0XECT
LR0JECT	LR0PECT	LR0QECT	LR0VECT	LR0XECT
PROJECT ※	PROLECT	PROQECT	PROVECT ※	PROXECT
QR0JECT	QR0LECT	QR0PECT	QR0VECT	QR0XECT
VR0JECT	VR0LECT	VR0PECT	VR0QECT	VR0XECT
XR0JECT	XR0LECT	XR0PECT	XR0QECT	XR0VECT

語として一致した2候補

PROJECT SOURCE (02=P、14=J)

PROJECTには計画・研究事業の意味がある。本文の「研究施設を建てた」という内容、本の表題としての「計画名」という読みとつながるため、こちらを優先する。これは文脈からの推定である。

PROVECT SOURCE (02=P、14=V)

PROVECTには古い形容詞として「成熟した・年齢が進んだ」の語義がある。言語学の動詞用法もある。「成熟した源」と読める可能性を残すが、この本の内容との結びつきは弱い。古語だからという理由だけでは除外しない。

見かけが似ていても条件を満たさない候補

PROTECTはID 14にもTを割り当てる必要があり、末尾のT=ID 10と一対一条件が衝突する。PROSECT・PROFECTも既採用のS=ID 15・F=ID 19と衝突する。

調べた範囲と、残る制約

macOSのweb2・web2aを大文字化して完全一致で照合。web2ではPROJECT・PROVECT、web2aでは該当なし。残り28候補はこの2一覧に未収録であり、「存在しない語」と証明したわけではない。固有名詞・造語・別言語・別区切りは網羅していない。

2候補に共通するPを暫定採用し、採用済みは21文字。ID 14はJ/Vのまま保留し、次の断片に現れる別の語で再確認する。

照合日：2026-09-16。辞書ファイルの行数・SHA-256、転記座標、前提はcover-investigation.json、全割り当てはcover-candidates.csvに保存。

語義の参照：[Merriam-Webster: project](#) / [Merriam-Webster: source](#) / [Wiktionary: provect \(古い形容詞\)](#) / [Merriam-Webster: provect \(言語学の動詞\)](#)

次の断片へ：調査の進め方を残す

今回までの調査

- ① 26字形に固定IDを付け、点・線・接続を原画像と照合してSVG化。② 最初の既知表を使わない版では、ガイドの訳文から13文字、本文から7文字を導いた。各Stepの開始時点の対応を保存した。
- ③ M・Tは、その時点の未使用12文字から132通りを列举し、3語を同時に比較。④ 全文をローマ字・ひらがな・漢字混じりで保存。「ならば」は文語として保持し、DOURIの表記整理とZの単独例を注記した。
- ⑤ 表紙の既知部分は？RO？ECT / SOURCE。未使用6文字から30通りを作り、辞書照合でPROVECTも発見した。PROJECTを優先するが、J/Vは未確定として次回へ持ち越す。

新しいページ・断片を受け取ったら

1. 原画像・入手日・ページ位置・向き・SHA-256を保存する。補正した画像は別ファイルとし、原画像へ戻れるようにする。
2. その時点の対応表を固定して保存する。ID→採用文字と、候補だけの文字を分ける。現在の続きはcurrent-alphabet.jsonから始める。
3. 文字数・行分け・反復を記録して固定IDへ転記する。薄い線や点は候補IDと切り抜き座標を残し、意味に合わせて補わない。
4. 既知文字だけを置換し、未判明は「？」のままにする。日本語ローマ字・英語・固有名詞を検討し、古い語法も調べる。
5. 未使用文字の重複しない割り当てを全列举する。複数の語・同じIDの全出現を同時に評価し、候補・除外理由・辞書の範囲を記録する。
6. 機械的な一致と文脈判断を分ける。優先する仮説、暫定採用、保留、反証された案を保存する。後の情報を前のStepへ混入させない。
7. 原文の綴り・読み・解釈・参考訳を分けて追記する。対応の対一性と出力を検査し、画像と最終PDFを目視確認する。

次回の重点：ID 14がJかVか、ID 02=Pの独立した例、Z=ID 17の再出現。未割り当てIDは05・08・14・24・25、文字はJ・L・Q・V・X。

詳細と記録テンプレート：investigation-process.md / 文体の判断：decoding-guidelines.md。確認する前に新しい字形を既存のどれかへ無理に当てはめない。

付録：アルファベット一覧 A～Z

この資料で導いた対応をアルファベット順に掲載。採用21文字（Pは暫定）・未割り当て5文字。小数字は字形の固定ID。

A  ID 20 訳で照合	B  ID 01 本文で照合	C  ID 07 訳で照合	D  ID 22 本文で照合	E  ID 13 訳で照合	F  ID 19 訳で照合	G  ID 11 訳で照合
H  ID 09 訳で照合	I  ID 23 訳で照合	J  ID 未確定 未判明	K  ID 04 訳で照合	L  ID 未確定 未判明	M  ID 18 本文で照合	N  ID 03 訳で照合
O  ID 21 訳で照合	P  ID 02 表紙・暫定 ※	Q  ID 未確定 未判明	R  ID 12 訳で照合	S  ID 15 訳で照合	T  ID 10 本文で照合	U  ID 06 訳で照合
V  ID 未確定 未判明	W  ID 26 本文で照合	X  ID 未確定 未判明	Y  ID 16 本文で照合	Z  ID 17 単独例 ※		

訳で照合＝ガイドの日本語訳から導出 / 本文で照合＝複数の語・文脈で確認

- ※ Z（ID 17）はZEROの1例による推定。複数語で確認できた文字より根拠が弱く、別の文章で再確認する。
- ※ P（ID 02）は表紙の2候補に共通するため暫定採用。ID 14はJを優先するがVも残り、どちらにも割り当てない。

対応を割り当てていない5字形

				
05	08	14	24	25
?	?	?	?	?

候補はJ・L・Q・V・X。ID 5・8・14・24・25との個別対応は保留。上のアルファベットと下の字形の並び順は対応を示さない。