

EXPLORING CONTINUOUS TENSEGRITIES

by

EDWARD BRUCE “TED” ASHTON

(Under the direction of Dr. Jason Cantarella)

ABSTRACT

A discrete tensegrity framework can be thought of as a graph in Euclidean n -space where each edge is of one of three types: an edge with a fixed length (bar) or an edge with an upper (cable) or lower (strut) bound on its length. Roth and Whiteley, in their 1981 paper “Tensegrity Frameworks”, showed that in certain cases, the struts and cables can be replaced with bars when analyzing the framework for infinitesimal rigidity. In that case we call the tensegrity *bar equivalent*. In specific, they showed that if there exists a set of positive weights, called a positive *stress*, on the edges such that the weighted sum of the edge vectors is zero at every vertex, then the tensegrity is bar equivalent.

In this paper we consider an extended version of the tensegrity framework in which the vertex set is a (possibly infinite) set of points in Euclidean n -space and the edgeset is a compact set of unordered pairs of vertices. These are called *continuous tensegrities*. We show that if a continuous tensegrity has a strictly positive stress, it is bar equivalent and that it has a semipositive stress if and only if it is partially bar equivalent. We also show that if a tensegrity is *minimally bar equivalent* (it is bar equivalent but removing any open set of edges makes it no longer so), then it has a strictly positive stress.

In particular, we examine the case where the vertices form a rectifiable curve and the possible motions of the curve are limited to local isometries of it. Our methods provide an attractive proof of the following result: There is no locally arclength preserving motion of a circle that increases any antipodal distance without decreasing some other one.

INDEX WORDS: Bar Equivalence, Dissertation, Framework, Infinitesimal Motion, Infinitesimal Rigidity, Stress, System of Inequalities, Tensegrity, Theorem of the Alternative

EXPLORING CONTINUOUS TENSEGRITIES

by

EDWARD BRUCE “TED” ASHTON

B.S.E., Walla Walla College, 1993

B.S., Southern Adventist University, 1999

A Dissertation Submitted to the Graduate Faculty
of The University of Georgia in Partial Fulfillment
of the
Requirements for the Degree
DOCTOR OF PHILOSOPHY

ATHENS, GEORGIA

2007

© 2007

Edward Bruce “Ted” Ashton

All Rights Reserved

EXPLORING CONTINUOUS TENSEGRITIES

by

EDWARD BRUCE “TED” ASHTON

Approved:

Major Professor: Dr. Jason Cantarella

Committee: Dr. Malcolm R. Adams
Dr. Edward A. Azoff
Dr. Joseph H. G. Fu
Dr. Robert Varley

Electronic Version Approved:

Dr. Maureen Grasso
Dean of the Graduate School
The University of Georgia
May 2007

DEDICATION

To my wife and best friend,

Heidi

ACKNOWLEDGMENTS

It is rare in our lives that we reach a point of completion, but the end of a doctorate is certainly such a point. As I pause to contemplate, I find that there are many to whom I owe a debt of gratitude.

Certainly the Chief Contributor is the One who hid Tensegrity within the cell long before we discovered it and who gives us strength even when the stresses don't appear to be strictly positive. My thanks is ever to Him.

Among the people who have helped to bring this day about, my family must head the list. My wife has given of herself in myriad ways. She has made a home for us all and sacrificed much in the pursuit of this degree. Her worth is truly far above rubies.

I'm grateful, also, for the love and support of our children, Jo and Jim, as well as our "other children", Kristen, Jill and Gene. I have been immensely grateful for the encouragement of our ten siblings and particularly of our loving parents, who have supported us emotionally, spiritually and financially.

Many are the friends and family who have provided aid and comfort. In particular, I'm grateful to Aunt Rilla for giving us true vacation times and to Judy deLay, the Nelsons and the Crosses all for their hospitality. The Athens Seventh-day Adventist Church has been a true church home for us. They welcomed us with open arms when we arrived and have continued to love and support us ever since. When I think of our friends and extended family around the globe, I am truly amazed at how blessed we are.

In the mathematical realm, my advisor, Dr. Jason Cantarella, leads the pack. He took me on as an apprentice long before I started my dissertation and poured many hours, sometimes entire days, into this work. He has taught me and advised me and

many times provided just the clue I needed. I am grateful for his depth of knowledge of the field and his patience in explaining it to me. I am doubly honored to have an advisor whom I can also call my friend.

Untold numbers of people have been involved in getting me to where I am now, far more than I can mention here. But there are a few who deserve particular notice. The members of my committee, Drs. Malcolm Adams, Edward Azoff, Joseph Fu and Robert Varley all gave time to this work. Dr. Adams provided many valuable editorial notes. Dr. Varley has answered numerous questions, both mathematical and otherwise. Dr. Azoff and Dr. Fu both spent countless hours answering questions and teaching me such things as Functional Analysis and Linear Algebra. Their sage advice and deep knowledge were the well from which I repeatedly drew in my efforts. I have had many fine professors here and I thank them all, but one more who deserves a mention is one from whom I never took a class. Dr. Ted Shifrin spent his time and energy going over exams I had designed and giving me advice on teaching and I am a better teacher because of it.

However, he was not the first to encourage me in teaching. My fifth grade teacher, Mrs. Barbara Stanaway, gave me my first taste of math tutoring and I have loved it ever since. Mr. Morford, one of my high school math teachers, was the first to put me in front of a class and has always encouraged me in my pursuit of mathematics. Years later he again trusted me with his class and has followed my progress with interest.

Many others have not only taught me but also cared about me as a student: Mr. Magi, who gave me my first taste of Geometry and Dr. Moore, who took me on from where he left off; Dr. Richert, who showed me that Mathematics is a performance art; Drs. Kuhlman and Hefferlin, the world-class physicists who invited me to be a physics major and kept on investing time in me even when I turned them down; Dr. Haluska, who taught me to write and Dr. Sylvia Nosworthy, who taught me (amongst many things) that my writing needs more fluff; Dr. Masden, who not only demanded

great things of his students, but always went the distance with us; Dr. Barnett, who opened to me the world of simulation and modeling; Dr. Rob Frohne who showed me the miracle of feedback and control; Ralph Stirling, who was not only a fantastic resource for an engineering student, but also a friend; Dr. Cross and his family, who not only provided good education, but opened their home to us; Dr. Wiggins, who started me into Numerical Analysis; Dr. Tim Tiffin, who gave me my first serious taste of mathematical research; and Dr. Tom Thompson, who taught and encouraged me not only in Mathematics, but also in learning $\text{\LaTeX}$, knowledge I have put to good use.

My thanks goes to these and the other professors I have had who have chosen to give their lives in service and who care deeply about their students. In particular, I am grateful to Drs. Wiggins, Thompson, Richert and Moore for taking a lonely graduate student under their wings at the National Math Meetings and providing for Sabbath away from home.

My thanks also goes to John Pardon, for an illuminating conversation about non-convex curves, to Branko Grünbaum for providing a copy of “Lost Mathematics” and to Kenneth Snelson for his gracious reply to this graduate student’s email about his sculptures. I appreciate also the camaraderie of my fellow graduate students: Jeremy Francisco, Becca Murphy, Amy Kelly, Valerie Hower, Alan Thomas, Xander Faber, Bryan Gallant, David and Alicia Beckworth and my officemates Kenny Little and Chad Mullikin and also of the (then) postdocs Aaron Abrams and Nancy Wrinkle. As I look at the piles of library books on my desk, I thank the people of the University of Georgia Science Library. Rarely have I needed a resource that they did not provide.

Finally, I think of Grandmother, who, eighty years ago, was a math major herself. She would be proud.

TABLE OF CONTENTS

	Page
ACKNOWLEDGMENTS	v
LIST OF FIGURES	x
CHAPTER	
1 INTRODUCTION: TENSEGRITY	1
1.1 A HISTORY	1
1.2 TENSEGRITY IN MATHEMATICS	3
1.3 PRESTRESSED CONCRETE	4
1.4 MATHEMATICAL CONTEXT	5
1.5 A ROAD MAP	8
2 THE ROTH-WHITELEY THEOREM	14
2.1 SETTING UP A ROTH & WHITELEY EXAMPLE	14
2.2 WORKING THE CROSSED-SQUARE EXAMPLE	24
2.3 THE MECHANICS	26
2.4 KEY LEMMA AND THEOREM	34
3 MAIN THEOREMS	37
3.1 UNDER THE HOOD: THEOREMS OF THE ALTERNATIVE	37
3.2 TURNING SETS INTO SPACES: $\mathcal{V}$ AND $\mathcal{E}$	39
3.3 $C(\mathcal{E})$ AND $C^*(\mathcal{E})$	52
3.4 $\text{VF}(\mathcal{V})$ AND Y	60
3.5 FIRST MAIN THEOREM	64

3.6	SECOND MAIN THEOREM	68
3.7	TWO EXAMPLES	73
3.8	MINIMAL $\mathcal{X}$ -BAR EQUIVALENCE	86
3.9	COVERED TENSEGRITIES	89
4	EXAMPLES	93
4.1	RECTANGLE	93
4.2	ON A CIRCLE	97
4.3	ON A SPHERE	99
4.4	ON A GENERAL CURVE	99
4.5	UNDERSTANDING A QUADRILATERAL	100
4.6	OTHER FRACTIONS	103
4.7	NON-CONVEXITY	105
4.8	AFFINE TRANSFORMATIONS	108
4.9	A CORNER CASE	115
5	CONCLUSION AND FUTURE DIRECTIONS	119
	BIBLIOGRAPHY	123
	APPENDIX	
A	EUCLIDEAN MOTIONS: INFINITESIMAL RIGID MOTIONS OF SPACE	129
A.1	THE SPECIAL ORTHOGONAL GROUP: SO_n AND $\mathfrak{so}_n$	129
A.2	THE SPECIAL EUCLIDEAN GROUP: SE_n AND $\mathfrak{se}_n$	133
B	MOTZKIN'S THEOREM	137
B.1	PREPARING TO USE MOTZKIN'S THEOREM	137
B.2	MOTZKIN'S THEOREM APPLIED	143
B.3	MOTZKIN ON THE OTHER HAND	145
B.4	STRESSES AS VARIATIONS	146
	INDEX	148

LIST OF FIGURES

1.1.1	A figure from Kenneth Snelson’s tensegrity patent.	2
1.3.1	A concrete slab supported at the ends.	4
1.3.2	Forces in a prestressed concrete slab.	4
1.4.1	A lemma of Cauchy reproven by Connelly (1982).	6
1.4.2	The first three regular Grünbaum polygons.	7
1.5.1	The circle-of-struts example of Section 3.7.	10
1.5.2	Works of Kenneth Snelson.	13
2.1.1	The crossed-square example with its bar framework.	15
2.1.2	Edges bearing negative loads.	19
2.1.3	Evidence that $I(p) \neq \bar{T}(p)$ in general.	20
2.1.4	A two-bar tensegrity with a motion that is in $\bar{T}(p)$ but not in $T(p)$	20
2.2.1	The crossed-square example with its edges numbered.	24
2.2.2	Every motion of the crossed-square bar framework is Euclidean.	26
2.3.1	Convex set generated by $(0, 0)$, $(1, 0)$ and the direction $(1, 1)$	30
3.1.1	A Theorem of the Alternative in $\mathbb{R}^2$	37
3.2.1	A tensegrity with noncompact $\mathcal{S}$ and $\mathcal{C}$ but compact $\mathcal{S} \cup \mathcal{C}$	42
3.2.2	Open sets in $\mathcal{T}_{p(\gamma)}$ are open in $\mathcal{T}_\gamma$	44
3.2.3	A tensegrity with “free vertices”.	47
3.2.4	The stadium-curve tensegrity.	49
3.2.5	The edgeset of the stadium-curve tensegrity	50
3.2.6	The tensegrity of Figure 3.2.1 “compactified”.	51
3.5.1	A tensegrity with a semipositive stress but no strictly positive one.	67
3.5.2	A tensegrity with no semipositive stress.	68

3.6.1	A non-minimal crossed square.	70
3.7.1	The circle-of-struts example.	74
3.7.2	The almost-half-a-circle-of-struts example.	74
3.7.3	The almost-half-a-circle-of-struts with the “good” motion V_g	84
3.8.1	Two squares of struts and a semipositive motion.	87
3.8.2	The $\mathcal{X}$ -bar equivalent subtensegrities of the squares of struts.	88
3.8.3	The non-minimal crossed square with its minimal subtensegrities.	89
3.9.1	A cylinder filled with antipodal struts.	91
4.1.1	The rectangle example.	94
4.1.2	The edges connecting to an internal vertex in the rectangle.	95
4.2.1	Two examples of the on-a-circle example.	97
4.2.2	The angles involved in the on-a-circle example.	98
4.3.1	The octahedron tensegrity.	99
4.3.2	Octahedra cover the spherical tensegrity.	99
4.4.1	The on-a-circle example with more general curves.	100
4.5.1	A general quadrilateral tensegrity.	100
4.6.1	A more general hexagon.	103
4.6.2	Two struts with their associated cables.	104
4.6.3	The hexagon with a motion.	104
4.6.4	Two regular figures that arise with this construction.	105
4.6.5	A bar equivalent tensegrity that is a nonregular hexagon.	105
4.6.6	Another in the family of nonregular, bar equivalent hexagons.	106
4.7.1	Two nonconvex examples.	107
4.7.2	“Double skip” and bar equivalence.	107
4.7.3	Vertex curve and strut that cannot be balanced.	108
4.8.1	The crossed-square, original and transformed.	108
4.8.2	A motion that is in $\ker Y$ only before transformation.	111

4.8.3	$\mathcal{X}$ -bar equivalent before transformation but not after.	111
4.8.4	Two ellipse constructions.	114
4.9.1	Representatives of our family of triangular tensegrities.	115
4.9.2	Continuous tensegrity formed from triangle subtensegrities.	115
4.9.3	Balancing forces placed on a subtensegrity.	117
4.9.4	The forces present on one “arm” of the subtensegrity.	117
4.9.5	The tensegrity with the vector field described.	118
5.0.1	Moving the “convexifying” problem in to the continuous world. . .	120
A.1.1	Three different vector fields on a square made of cables.	129
A.2.1	$SE_n \not\cong SO_n \times \mathbb{R}^n$	134
A.2.2	Two vector fields in $T(p)$ on the crossed square.	136
B.1.1	The diagram for the Motzkin Transposition Theorem.	138
B.2.1	The Motzkin diagram and our matching situation.	143
B.3.1	Our new use of the Motzkin theorem.	145
B.4.1	The circle of struts with the vectors of $V_{d\theta}$	147

CHAPTER 1

INTRODUCTION: TENSEGRITY

1.1 A HISTORY

Stone is strong. Good quality stone can withstand amazing compressive forces. But in tension it is not nearly as strong. Similarly, brick and concrete are much stronger under compression than under tension (Table 1.1 gives some representative numbers). And wood, while fairly strong under tension parallel to the grain, can't compare to the durability of stone. So for millenia, buildings were designed to be continuous frameworks of members in compression (Pugh, 1976; Girvin, 1944; Forest Products Laboratory, 1999).

Table 1.1: Ultimate stresses for certain materials. Numbers have been converted to consistent units and rounded. Bending stress given for wood is actually modulus of rupture. Values for bone are from (Martin et al., 1998) and are given for interest, since tensegrity theory is used to understand the musculoskeletal system (see, for example, Chen and Ingber (1999)). Values marked ^a are from Forest Products Laboratory (1999); those marked ^b are from Girvin (1944), and those marked ^c come from MatWeb (2007).

Material	Average Ultimate Stress (lb/in ²)		
	Compression	Tension	Bending
Granite	14,000–45,000 ^c	1,020–3,630 ^c	1,600 ^b
Brick	10,000 ^b	200 ^b	600 ^b
Steel	60,000–72,000 ^b	60,000–72,000 ^b	60,000–72,000 ^b
Concrete	2,030–10,200 ^c	—	—
Sugar Maple (para. to grain)	8,000 ^a	15,700 ^a	16,000 ^a
Pin Oak (para. to grain)	7,000 ^a	16,300 ^a	14,000 ^a
Human bone (longitudinal)	28,000	19,000	30,000

The middle of the nineteenth century saw a number of major advances in the creation of steel, including the discovery of its microstructure and the resulting development of the science of metallurgy (Grun, 1979). This provided a building material with excellent tensile strength (see, for example, Shackelford and Alexander (2001)).

In 1965, Kenneth D. Snelson patented what he called “Continuous Tension, Discontinuous Compression Structures” (Snelson, 1965) (see Figure 1.1.1 and Figure 1.5.2). His new creations had cables under tension running throughout and then, at intervals along the cables, steel struts under compression.

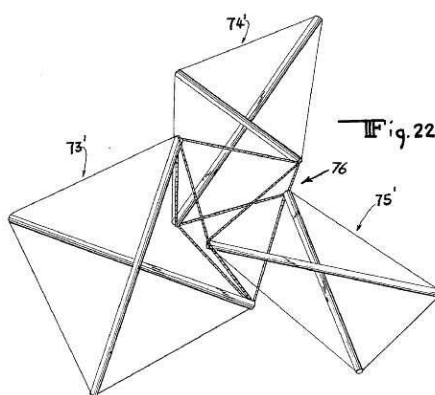


Figure 1.1.1: A figure from patent 3,169,611 by Kenneth Snelson (1965).

Buckminster Fuller recognized in Snelson’s idea new possibilities in Architecture that would allow for much lighter, more efficient structures. He termed the concept “Tensional Integrity” or “Tensegrity”. Since then, Tensegrity has been used in fields as diverse as Architecture, Cellular Biology, Dairy Science, Dance, Ornithology, Robot Kinematics and the studies of sleep disorders and the musculoskeletal system. (Pugh, 1976; Eaves, 2004; Volokh et al.; Farrell et al.; Franklin, 1996; Jeong et al., Jan. 2007; Maina, 2007; Lenarčič and Gallettti, 2004; Roth, Apr. 2005; Chen and Ingber, 1999).

1.2 TENSEGRITY IN MATHEMATICS

In 1981, Ben Roth and Walter Whiteley wrote “Tensegrity Frameworks” (Roth and Whiteley, 1981) in which they showed how to extend the tools of rigidity analysis on bar frameworks (n -dimensional graphs in which the edge lengths are fixed) to certain tensegrity frameworks. This was one of the earliest works treating tensegrities in the mathematical realm and will be a central motivator in what we do here.

For us, a *tensegrity* will consist of four sets and a map. First, there will be a set $\mathcal{V}$ of vertices. We will want these to be points in Euclidean n -space, but we’ll also want to be able to think about these points moving, so we’ll take $\mathcal{V}$ to be an abstract set accompanied by a 1-1 continuous map $p: \mathcal{V} \rightarrow \mathbb{R}^n$ that places the elements of $\mathcal{V}$ in space. Next, there are three sets of edges connecting various pairs of vertices, with the understanding that at most one edge connects any given pair.

These edge sets are the set of *struts*, $\mathbf{S}$, the set of *cables*, $\mathbf{C}$, and the set of *bars*, $\mathbf{B}$. Two vertices connected by a strut are never allowed to move any closer together than they already are, though they may move farther apart. Two vertices connected by a cable cannot move farther apart but may move nearer. And two vertices connected by a bar must stay exactly as far apart as they currently are.

If the only edges a tensegrity has are bars, we may call the tensegrity a *bar framework*.

To date, tensegrities have been finite affairs, with a finite number of finite elements. But the time has come to look outside that realm. We will take $\mathcal{V}$ to be an arbitrary set. That will open the door to having continuous families of edges and so we will call our new constructions “continuous tensegrities”.

1.3 PRESTRESSED CONCRETE

Concrete, like stone, is fairly strong in compression and not nearly so strong in tension (Girvin, 1944, p. 213). Unfortunately, a large slab of concrete supported only at the ends experiences a bending load that puts the lower portion of the concrete in tension (see Figure 1.3.1) (Girvin, 1944, pp. 81–82).

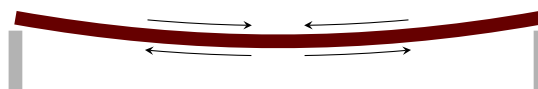


Figure 1.3.1: A concrete beam or slab supported only at the ends will experience tension below and compression above.

Prestressed concrete is a method of dealing with that problem. Steel rods or cables are installed under tension in the concrete. Their tension puts the concrete in compression (see Figure 1.3.2). Now, up to a certain load, the concrete remains under compression, functioning in its ideal realm (Wikipedia, 2007a).

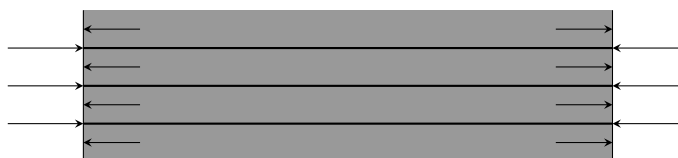


Figure 1.3.2: Forces in a prestressed concrete slab. The tension in the steel rods pulls inward on the faceplates, while the compression in the concrete pushes outward, placing them at equilibrium. As a tensegrity, the concrete acts like a strut and the rods like a cable, while the faceplates serve as vertices. This diagram is of post-tensioned prestressed concrete, where the tension is added to the rods after the concrete has hardened. In pre-tensioning, the concrete is poured around and bonds to rods that are already in tension. The forces then act all along the rods instead of only at the ends.

The prestressed concrete is a tensegrity framework with a *stress*, that is, a set of compressions on its struts (the concrete) and tensions on its cables (the rods) that gives a net zero force at every vertex (faceplate).

A main result in “Tensegrity Frameworks” (Roth and Whiteley, 1981) is that if (and only if) a tensegrity can be stressed in this fashion, then it functions no longer

as if it had cables and struts but rather as if all of its elements were bars. Moving this theorem into the realm of continuous tensegrities is our primary goal.

1.4 MATHEMATICAL CONTEXT

Let's look, for a moment, at a few places where this type of theorem has been used to good effect.

In recent years, two teams have settled the Carpenter's Rule Problem, which asks, given a polygonal arc with fixed-length edges in the plane, whether it is always possible to straighten the arc without it intersecting itself during the process (Connelly et al., 2003; Rote et al., 2003).

Connelly, Demaine and Rote (2003) not only settle the question in the affirmative but also answer the related question of whether a closed, simple polygon with fixed-length edges in the plane can be made convex through some motion that, again, avoids self intersections. In fact, they show that there is a flow of the polygon or arc that accomplishes the purpose and during which all of the distances between nonadjacent vertices are strictly increasing.

Their work uses a theorem very much like Roth & Whiteley's. They turn the polygon into a tensegrity and show that if it has no "weak stress", then it has a strictly expansive motion. John Pardon (private communication) has made significant progress in extending this to smooth curves, showing that for a simple, closed, nonconvex curve, there exists a homotopy that takes it to a convex curve, during which all of the self-distances on the curve are nondecreasing.

In our second example, Connelly applies Roth & Whiteley's theorem to provide an additional proof for the following lemma, which he credits to Cauchy (Connelly, 1982) (see Figure 1.4.1).

Lemma 1.4.1. *If, in a convex planar or spherical polygon $ABCDEF$, all the sides $AB, BC, CD, \dots, FG$, with the exception of only AG , are assumed invariant, one*

may increase or decrease simultaneously the angles B, C, D, E, F included between these same sides; the variable side AG increases in the first case, and decreases in the second.

Proof. See (Connelly, 1982, p. 30). □

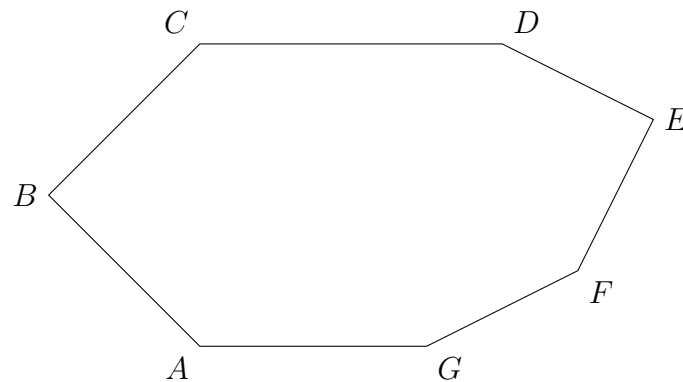


Figure 1.4.1: Increasing (resp. decreasing) angles B through F lengthens (resp. shortens) side AG . A lemma of Cauchy reproven by Connelly (1982).

In his discussion, Connelly notes that this is the polygonal equivalent of Schur’s Theorem, which says that if we have a C^2 convex plane arc and we decrease its curvature at some place or places and do not increase it anywhere, then the distance between its ends increases (see, for example, do Carmo (1976, p. 406)).

Our final example comes from Roth & Whiteley themselves. Grünbaum and Shephard, in their “Lectures on Lost Mathematics” (Grünbaum and Shephard) introduce a class of tensegrities that Roth & Whiteley call the “Grünbaum polygons”. These are convex, planar (not necessarily regular) polygons made of struts that have cables joining a select vertex to all nonadjacent vertices and one further cable joining the two vertices adjacent to the chosen one (see Figure 1.4.2 on the next page).

Roth and Whiteley (1981, p. 437) show that the Grünbaum polygons are infinitesimally rigid in $\mathbb{R}^2$. As the number of edges in the Grünbaum polygons increases, they seem to approach a smooth convex curve. So it seems not unreasonable that the limit of such a family would be ...

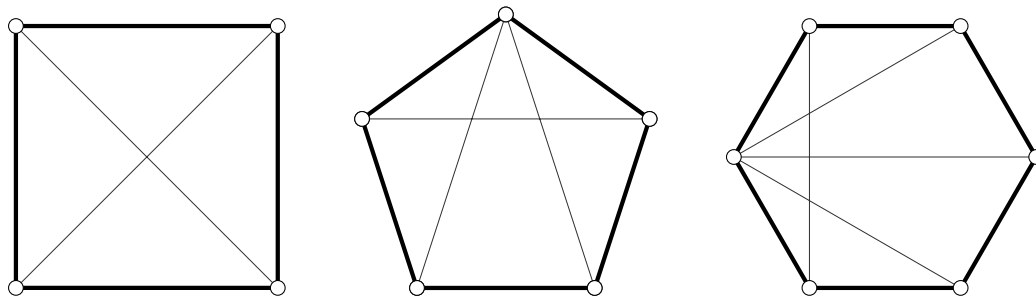


Figure 1.4.2: The first three regular Grünbaum polygons, introduced in (Grünbaum and Shephard). Struts are represented by thick lines, cables by thin ones, vertices by circles.

... a convex curve in $\mathbb{R}^2$ that is not allowed to shrink locally to first order (but may stretch) and ...

... which has a continuous family of upper bounds on the distances between one distinguished point on the curves and all the others and ...

... some sort of curvature requirement at that point,

and that this limiting curve would be infinitesimally rigid.

Here we have three situations in which Tensegrity Theory has been used to solve problems in Discrete Geometry, problems that have analogs in Differential Geometry. Furthermore Cantarella, Fu, Kusner, Sullivan and Wrinkle (2006) borrow terms and ideas from Tensegrity Theory to answer questions in Thick Knot Theory, distinctly a Differential Geometry field.

One can hardly avoid the feeling that these and other problems would benefit from having tensegrity tools brought from the discrete realm into the continuous, and we begin that project with this work.

1.5 A ROAD MAP

Here in [Chapter 1](#), we’ve given a history, both structural and mathematical, placing this work in context. In [Chapter 2](#) we build up the background for the main theorems. [Section 2.1](#) uses the setting of a specific example tensegrity to introduce the notation and terms we will need in the rest of the paper. For example, we here meet the *variations* (Definition [2.1.3](#) on page 17), the possible ways of moving the vertices; the (infinitesimal) *motions* (Definition [2.1.7](#) on page 18), the variations that respect all of the edge constraints; and the *loads* (Definition [2.1.8](#) on page 19), the effects of variations on the edges of the tensegrity. We’ll discover that there are *strictly positive* motions, which increase the lengths of all struts and decrease the lengths of all cables and the *semipositive* ones, which increase the length of at least one strut or decrease the length of at least one cable.

We also find out about *infinitesimally rigid* tensegrities, whose only motions are rigid motions of space and about *bar equivalent* tensegrities, which act as if they were made entirely of bars, and *partially bar equivalent* tensegrities, which have bar equivalent pieces (Definitions [2.1.9](#) on page 20 and [2.1.12](#) on page 21).

We will take, as the variations of interest, a subspace, $\mathcal{X}$, of the space of variations, called the *design variations* (Definition [2.1.6](#) on page 18). If $\mathcal{X}$ doesn’t contain all of the variations, we’ll use the terms *$\mathcal{X}$ -infinitesimally rigid*, *$\mathcal{X}$ -bar equivalent* and *partially $\mathcal{X}$ -bar equivalent*.

Finally, we’ll meet the *stresses* (Definition [2.1.16](#) on page 22), sets of tensions on the cables and compressions on the struts that, like those in the prestressed concrete, put the tensegrity in equilibrium. These also come in strictly positive (positive values on all the edges) and semipositive (a positive value on at least one edge).

In [Section 2.2](#) we work the example we have set up in [Section 2.1](#). [Section 2.3](#) covers the mathematical background needed for the proofs in [Section 2.4](#). Then in

Section 2.4, we state and prove Roth & Whiteley’s key lemma, which we reformulate as follows:

Lemma 2.4.3 (p. 35). *Let $G(p)$ be a (finite) tensegrity framework with rigidity matrix Y . Then $G(p)$ is bar equivalent if and only if $G(p)$ has a strictly positive stress.*

Chapter 3 holds our main results. In Section 3.1 we talk about a class of theorems called “Theorems of the Alternative”. Our main theorems resemble these. In fact, in Appendix B, we will show how to prove a weaker version of our first main theorem using a Theorem of the Alternative by Theodore S. Motzkin.

Next we launch into the task of moving the things we have defined in Chapter 2 into the continuous realm. Section 3.2 builds a metric and a topology on the vertex set $\mathcal{V}$ and the *edgeset* $\mathcal{E}$. Section 3.3 takes on $C(\mathcal{E})$ and $C^*(\mathcal{E})$ (the set of continuous functions on $\mathcal{E}$ and its topological dual), giving them norms and topologies and then exploring their natures in some depth. In Section 3.4 we address $\text{VF}(\mathcal{V})$ and Y .

Then, in Section 3.5 we reach our first main theorem:

Theorem 3.5.2 (p. 65). *The tensegrity $G(p)$ is partially $\mathcal{X}$ -bar equivalent if and only if $G(p)$ has a semipositive stress.*

We explore the theorem with a couple of examples. Then we move, in Section 3.6, to our second main theorem:

Theorem 3.6.1 (p. 68). *If a tensegrity has a strictly positive stress, then it is $\mathcal{X}$ -bar equivalent.*

But we only have one direction of the desired theorem and we do some looking at why the other direction is difficult.

Section 3.7 is an excursion into example land. We take on a couple of examples of tensegrities in which $\mathcal{X}$ is not the whole of $\text{VF}(\mathcal{V})$.

We discover that the example called the “circle of struts” is bar equivalent. In this example, a circle of vertices is filled with antipodal struts (see [Figure 1.5.1](#)) and we only allow variations which are local isometries on the curve. When we apply the

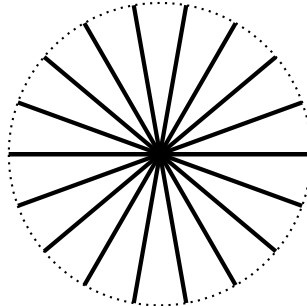


Figure 1.5.1: The circle-of-struts example of [Section 3.7](#).

first main theorem to this example, we get the following proposition:

Proposition 3.7.18 (p. 83). *There is no locally arclength preserving motion of a circle that increases all antipodal distances.*

When we apply the second main theorem, that becomes:

Proposition 3.7.19 (p. 83). *There is no locally arclength preserving motion of a circle that increases any antipodal distance without decreasing some other one.*

The circle acts like a water balloon. Pulling it out somewhere results in it being pulled in elsewhere.

In [Section 3.8](#) we again attempt to provide the “other half” of the second main theorem. By restricting our attention to tensegrities that are $\mathcal{X}$ -bar equivalent and that cease to be so upon the removal of any open set of edges, we get

Theorem 3.8.3 (p. 86). *If $G(p)$ is minimally $\mathcal{X}$ -bar equivalent, $G(p)$ has a strictly positive stress.*

After proving the theorem, we talk about what minimally $\mathcal{X}$ -bar equivalent tensegrities are like. Then, in [Section 3.9](#), we show that if a tensegrity is (at least mostly)

covered by subtensegrities which have strictly positive stresses, then the tensegrity itself has a strictly positive stress:

Theorem 3.9.2 (p. 90). *If $G(p)$ is countably covered by subtensegrities that have a strictly positive stress, then $G(p)$ has a strictly positive stress.*

We end by conjecturing that either

Conjecture 3.9.3 (p. 91). *Every $\mathcal{X}$ -bar equivalent continuous tensegrity is countably covered by minimally $\mathcal{X}$ -bar equivalent subtensegrities.*

or at least

Conjecture 3.9.4 (p. 92). *Every $\mathcal{X}$ -bar equivalent tensegrity has a strictly positive stress.*

Chapter 4 is full of examples of continuous tensegrities. In Section 4.1 we analyze a tensegrity that is not only bar equivalent but also infinitesimally rigid. Section 4.2 through Section 4.8 are a theme and variations. First we start by taking a circle of vertices and setting struts antipodally on it and cables connecting vertices a fixed skip around. That works so well that we move up one dimension and build the equivalent tensegrity on the sphere, thus providing ourselves with an example of a tensegrity whose vertex set is two-dimensional. Then we try to relax the circle into other vertex curves, with varying amounts of success. In particular, Section 4.8 shows that bar equivalence can survive affine transformations of the tensegrity.

Our last example (Section 4.9) is designed to remind us that the intuition we get about subspaces and closed convex cones from thinking about $\mathbb{R}^2$ and $\mathbb{R}^3$ may be faulty and to make a little clearer why showing that every bar equivalent tensegrity has a strictly positive stress is so difficult.

Finally, in Chapter 5, we wrap things up. We remind the reader of the situations we mentioned above where continuous tensegrities might bring some of the successes

of Discrete Geometry into Differential Geometry and we also point out some other places for further exploration. We close with two appendices. [Appendix A](#) covers rigid motions of space and the variations which are induced from them, and [Appendix B](#) proves a special-case version of the first main theorem using Motzkin's Transposition Theorem.

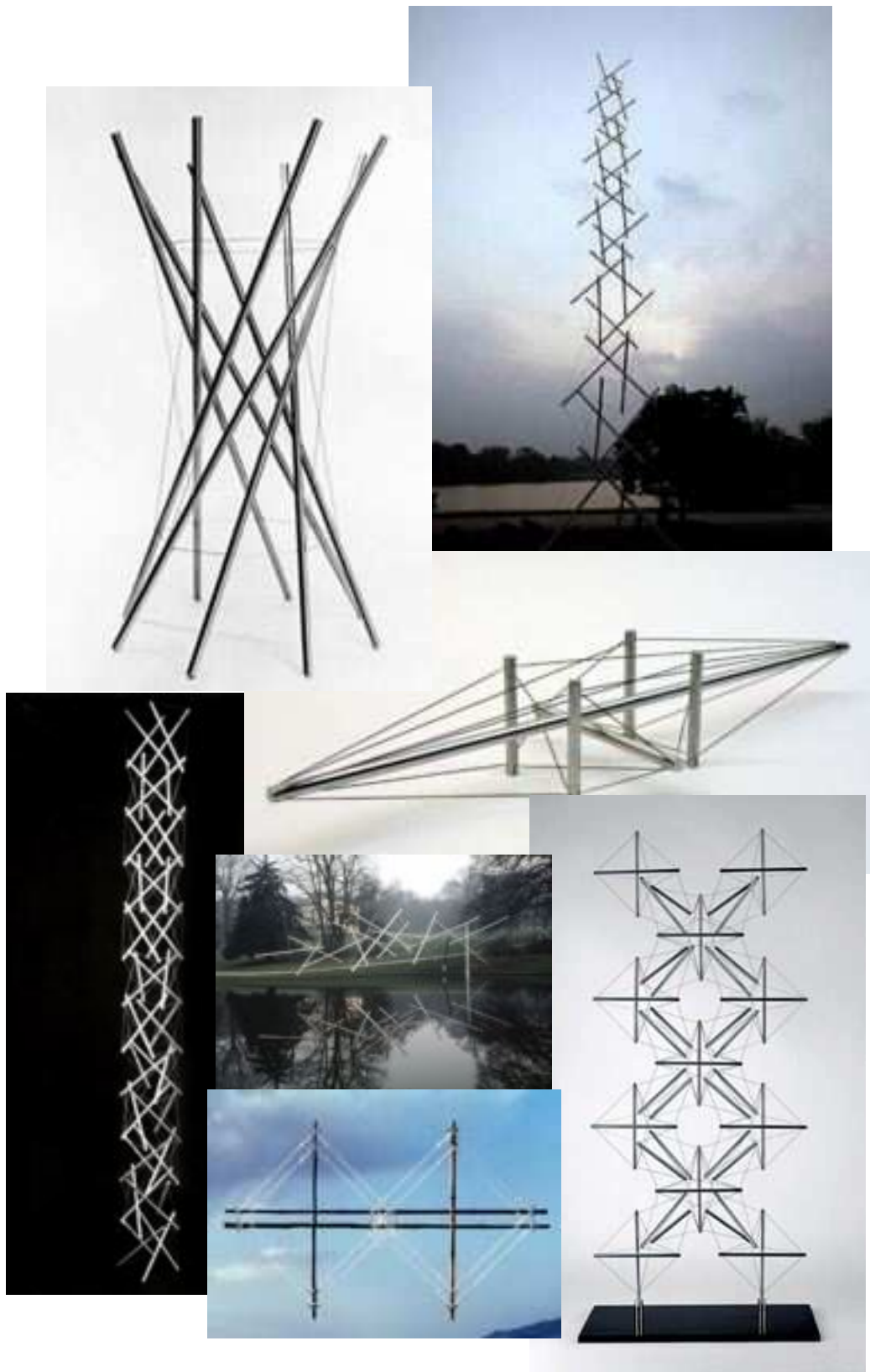


Figure 1.5.2: A collection of pictures of the works of Kenneth D. Snelson from his website www.kennethsnelson.net, used by permission.

CHAPTER 2

THE ROTH-WHITELEY THEOREM

2.1 SETTING UP A ROTH & WHITELEY EXAMPLE

We need to start by understanding Roth & Whiteley's original theorem. To that end, we'll first look at an example, and then work through their proof. However, we have another task to accomplish as well. There are numerous symbols and terms we'll want to have available when we get into the thick of the mathematics. As we work through the example, we'll define the terms and symbols, relating them to the example. Afterward we'll give a summary list of the symbols we have defined.

Let's get a little general notation out of the way first. In what follows, $\mathbf{0}$ will denote the origin of the various vector spaces we encounter. The symbol 0 will be reserved for the real additive identity. $\langle \cdot, \cdot \rangle$ will denote the inner product for any vector space. If the space is Euclidean n -space (denoted $\mathbb{R}^n$), we may use dot product notation (e.g. $v \cdot w$) for the same purpose.

If X and Y are two sets, then $X \setminus Y$ will denote the set of elements that are in X but not in Y . That is,

$$X \setminus Y := \{x \in X : x \notin Y\}.$$

We'll also want some idea of positivity for vectors and vector fields in $\mathbb{R}^n$.

Definition 2.1.1. A vector $v \in \mathbb{R}^n$ is said to be *strictly positive*, written $v > \mathbf{0}$, if every coordinate of v is strictly positive. v is said to be *nonnegative*, written $v \geq \mathbf{0}$,

if every coordinate of v is nonnegative. Finally, v is said to be *semipositive*, written $v \gneq \mathbf{0}$ if $v \geq \mathbf{0}$ and $v \neq \mathbf{0}$.

A vector field V is termed *strictly positive* (resp. *nonnegative*), also written $V > \mathbf{0}$ (resp. $V \geq \mathbf{0}$), if every value of V is strictly positive (resp. nonnegative) as a vector. V is termed *semipositive* (also written $V \gneq \mathbf{0}$) if $V \geq \mathbf{0}$ and V is not the “zero vector field”, that is, $V \neq \mathbf{0}$.

Now we can turn our attention to the example. Figure 2.1.1(a) shows a simple tensegrity. In our illustrations, struts will be shown as thick lines, cables as thin ones

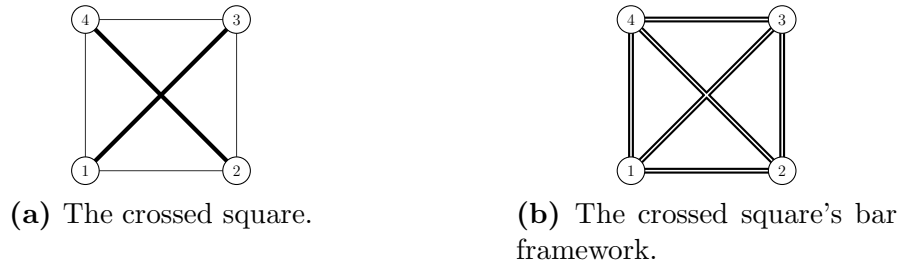


Figure 2.1.1: The crossed-square example with its bar framework. Note that the two struts (thick lines in (a)) and the two center bars (double lines crossing in (b)) pass through each other without touching.

and bars as double lines. Vertices may be shown as circles, as in Figure 2.1.1, or as dotted lines, as in Figure 1.5.1 on page 10.

Using the notation introduced on page 3, we see that the crossed square has

$$\begin{aligned} \mathcal{V} &= \{1, 2, 3, 4\} & \mathbf{S} &= \{\{1, 3\}, \{2, 4\}\} \\ \mathbf{B} &= \emptyset & \mathbf{C} &= \{\{1, 2\}, \{2, 3\}, \{3, 4\}, \{1, 4\}\} \end{aligned}$$

We need to place the elements of $\mathcal{V}$ in space, so we define $p: \mathcal{V} \rightarrow \mathbb{R}^2$ by

$$p(v) = \begin{cases} (0, 0), & v = 1 \\ (1, 0), & v = 2 \\ (1, 1), & v = 3 \\ (0, 1), & v = 4 \end{cases}$$

We'll call the tensegrity $G(p) = \{\mathcal{V}; \mathbf{S}, \mathbf{C}, \mathbf{B}; p\}$. [Figure 2.1.1\(b\)](#) shows the the same tensegrity but with all the struts and cables replaced by bars. The bar framework constructed from tensegrity $G(p)$ in this fashion will be called $\overline{G}(p)$.

Bars are a problem. They give us equations instead of inequalities. One might argue that we prefer equations to inequalities and perhaps that's true. However, in general, we can't turn inequalities into equalities. But we *can* turn an equation into a pair of inequalities by thinking of a bar as a strut-cable pair. That way the cable in the pair keeps the vertices from moving apart and the strut keeps them from moving together.

To make it clear that we are thinking of bars no longer as bars but as strut-cable pairs, we'll define the sets of “functional struts” and “functional cables”:

$$\mathcal{S} := \mathbf{S} \cup \mathbf{B} \qquad \mathcal{C} := \mathbf{C} \cup \mathbf{B}$$

and deliberately refer to the elements of $\mathcal{S}$ as “struts” and the elements of $\mathcal{C}$ as “cables”. We're ready now to give a name to the set of edges, which will contain no bars, only (functional) struts and cables:

Definition 2.1.2. The *edgeset* $\mathcal{E}$ is defined as a set by

$$\mathcal{E} = \mathcal{S} \sqcup \mathcal{C}$$

(where $\sqcup$ indicates disjoint union as opposed to the $\cup$ indicating union above). $\mathcal{E}$ is endowed with a topology which we will describe in [Section 3.2](#) on page 39 and required to be compact.

Here is a quick preview of how we make $\mathcal{V}$ and $\mathcal{E}$ into topological spaces. We induce a metric on $\mathcal{V}$ from $\mathbb{R}^n$ via p . We then take the maximum of two distances in $\mathcal{V}$ to give us a distance in $\mathcal{V} \times \mathcal{V}$. We use that metric on $\mathcal{V} \times \mathcal{V}$ to create a metric on the identification space $(\mathcal{V} \times \mathcal{V})/_\sim$ of unordered pairs of elements of $\mathcal{V}$. Finally we create a metric on $\mathcal{E}$ from the metric on $(\mathcal{V} \times \mathcal{V})/_\sim$.

Now that we have metrics for $\mathcal{V}$ and $\mathcal{E}$, we put the associated metric topologies on them. These turn out to be familiar topologies (Subsection 3.2.2 on page 43) and when the two sets are finite, both of the topologies are discrete (Proposition 3.2.12 on page 52). Of course, in that case both sets are compact simply because they are finite, but in the general case we require $\mathcal{E}$ to be compact and discover that $\mathcal{V}$ might as well be, too (Subsection 3.2.3 on page 47).

The possible motions of $G(p)$ can be described by putting vectors on the vertices. We'll call the space of all continuous vector fields on the vertices $\text{VF}(\mathcal{V})$ and note that in this case it is $\mathbb{R}^8$. Of course, only some of the elements of $\text{VF}(\mathcal{V})$ respect the edge constraints. We'll reserve the term “motion” to apply to those. If we want to talk generally about the elements of $\text{VF}(\mathcal{V})$, we'll refer to them as “variations” or just call them vector fields.

Definition 2.1.3. The *variations* of $G(p)$ are the continuous vector fields $V: \mathcal{V} \rightarrow \mathbb{R}^n$ and the space of variations is denoted $\text{VF}(\mathcal{V})$.

Now the variations are of different kinds. Some of the vector fields in $\text{VF}(\mathcal{V})$ can be induced from rigid motions of $\mathbb{R}^n$. They do things like translating $G(p)$ or rotating it, but they do not change its shape (see Appendix A on page 129 for more on rigid motions and the vector fields that are induced by them).

Definition 2.1.4. $T(p)$ is the set of all elements of $\text{VF}(\mathcal{V})$ that can be induced from rigid motions of $\mathbb{R}^n$. We'll term these *Euclidean motions*.

We need to know how a given vector field interacts with the constraints given by the edges. If the positions of the vertices vary differentiably with time and $\frac{\partial}{\partial t}p(v, t) = V(v)$ for a given vector field V , then V yields a vector of real numbers, which are the values of the derivatives of the lengths of the edges.

Here we have 6 edges, so that vector lies in $\mathbb{R}^6$, but since $\mathcal{E}$ has the discrete topology on it, we can (and will want later to) call it a “continuous function on $\mathcal{E}$ ”

and the space it inhabits $C(\mathcal{E})$. We'll call the functions we can generate this way *loads*. In a moment we'll want to talk about the nonnegative orthant of $C(\mathcal{E})$. We'll name that $C(\mathcal{E})^+$, and the nonpositive orthant we'll call $C(\mathcal{E})^-$.

Formally, loads are computed via the the *rigidity operator*:

Definition 2.1.5. $Y: \text{VF}(\mathcal{V}) \rightarrow C(\mathcal{E})$ is given by

$$(YV)(\{v_1, v_2\}) := \begin{cases} (V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2)), & \{v_1, v_2\} \in \mathcal{S} \\ -(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2)), & \{v_1, v_2\} \in \mathcal{C}. \end{cases} \quad (2.1)$$

The signs here were chosen so that $YV > \mathbf{0}$ when V expands struts and contracts cables. It would certainly be possible to do the work that follows using a Y which is defined the same for all parts of the edgeset. But we would do so at the cost of working in an orthant of $C(\mathcal{E})$ that has no natural name and of repeatedly saying “where $YV(e) \geq 0$ on struts and $YV(e) \leq 0$ on cables”.

In some cases we will not want to consider all possible vector fields. For example, we may have a curve of vertices and be interested only in variations that are local isometries of the curve. We will call the variations of interest *design variations*.

Definition 2.1.6. The *design variations* form a subspace $\mathcal{X} \subset \text{VF}(\mathcal{V})$ with the one requirement that the Euclidean motions be included in the design variations. That is, $T(p) \subset \mathcal{X}$.

With Y defined as above, for $V \in \mathcal{X}$, the statement $YV \in C(\mathcal{E})^+$ means that V is a variation on $\mathcal{V}$ that respects all of its constraints. As these are variations, they are really infinitesimal motions, but by a slight abuse of the language, we'll call them simply *motions*.

Definition 2.1.7. A *motion* is a variation $V \in \mathcal{X}$ such that $YV \in C(\mathcal{E})^+$. The set of all motions is denoted $I(p)$. A motion is termed *strictly positive* or *semipositive* as YV is strictly positive or semipositive.

So a strictly positive motion would increase the lengths of all struts and decrease the lengths of all cables, while a semipositive one would change the length of at least one edge.

And now we're prepared to formally define the term “load”.

Definition 2.1.8. The *loads* for a given tensegrity are the elements of the set $Y(\mathcal{X}) : = \{YV : V \in \mathcal{X}\}$. A load is *strictly positive* if $YV > \mathbf{0}$ and *semipositive* if $YV \geq \mathbf{0}$.

A quick note on sign. This definition of “load” is related to the idea of loads in Engineering, but it differs in sign. For example, an engineer would consider the edges in [Figure 2.1.2](#) to be bearing (positive) loads. We agree that they are bearing loads, but, for historical reasons, give them a negative sign.

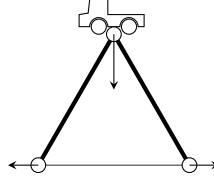


Figure 2.1.2: Edges bearing negative loads.

Some tensegrities possess motions that, while not being Euclidean motions, still don't change the length of any edges (see [Figure 2.1.4](#) on the following page for an example of such a tensegrity). Since such a motion changes no edge lengths, it respects the constraints not only of $G(p)$ but also of $\overline{G}(p)$.

So, similar to $I(p)$, we'll define $\overline{I}(p)$ to consist of the vector fields in $\mathcal{X}$ that respect the constraints of $\overline{G}(p)$. Because the constraints of $\overline{G}(p)$ are all equalities, any element $V \in \overline{I}(p)$ must give us $YV = \mathbf{0}$.¹ Clearly, $T(p) \subset \overline{I}(p) \subset I(p)$, but the reverse inclusions are not guaranteed, as [Figure 2.1.3](#) and [Figure 2.1.4](#) demonstrate.

¹In passing, we note that since Y is linear in V , any positive multiple of an element of $I(p)$ still maps by Y into $C(\mathcal{E})^+$ and the same is true of *convex* combinations of elements of $I(p)$, so $I(p)$ is a convex cone.

In contrast, all *linear* combinations of elements of $\overline{I}(p)$ must still map to $\mathbf{0}$, so $\overline{I}(p)$ is a subspace of $\text{VF}(\mathcal{V})$.



Figure 2.1.3: Evidence that $I(p) \neq \bar{I}(p)$ in general. The tensegrity on the left has a strut between the vertices, which allows them to move away from each other. The associated bar framework on the right does not allow that motion.

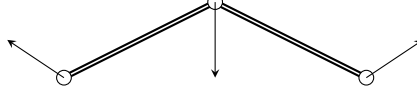


Figure 2.1.4: A two-bar tensegrity with a motion that is in $\bar{I}(p)$ but not in $T(p)$. Since this tensegrity is made only of bars, $I(p) = \bar{I}(p)$ naturally.

Now, for some tensegrities, we *do* have that $I(p) = \bar{I}(p)$ or $\bar{I}(p) = T(p)$ or even $I(p) = T(p)$. If $I(p) = T(p)$, then the only motions of $G(p)$ are the Euclidean motions (and similarly for $\bar{I}(p) = T(p)$). There are no motions that change the shape of $G(p)$ (or $\bar{G}(p)$). On the other hand, if $I(p) = \bar{I}(p)$, then the only motions of $G(p)$ are also motions of $\bar{G}(p)$. That is, the tensegrity acts as if it were made only of bars.

Definition 2.1.9. If $I(p) = T(p)$ (resp. $\bar{I}(p) = T(p)$), then we say that $G(p)$ (resp. $\bar{G}(p)$) is *infinitesimally rigid with respect to $\mathcal{X}$* or *$\mathcal{X}$ -infinitesimally rigid*.

If $I(p) = \bar{I}(p)$, we call $G(p)$ *bar equivalent with respect to $\mathcal{X}$* or *$\mathcal{X}$ -bar equivalent*.

In the case where $\mathcal{X} = \text{VF}(\mathcal{V})$ (as it is in Roth & Whiteley's work), we can drop the $\mathcal{X}$ and use the terms *infinitesimally rigid* and *bar equivalent*.

This leads to a quick proposition.

Proposition 2.1.10. *Any tensegrity that is infinitesimally rigid with respect to $\mathcal{X}$ is bar equivalent with respect to $\mathcal{X}$.*

Proof. $\mathcal{X}$ -infinitesimal rigidity means that $I(p) = T(p)$, but it is always true that $T(p) \subset \bar{I}(p) \subset I(p)$, so we must have $I(p) = \bar{I}(p)$, that is, that $G(p)$ is $\mathcal{X}$ -bar equivalent. □

There are three seemingly different ways of saying “bar equivalent”. They really all mean the same thing, however.

Proposition 2.1.11. *Let $G(p)$ be a tensegrity with motions $I(p)$, bar framework $\overline{G}(p)$ and $\overline{I}(p)$ the motions of $\overline{G}(p)$. Then the following are equivalent:*

1. $G(p)$ has no semipositive motions.
2. $Y(\mathcal{X}) \cap C(\mathcal{E})^+ = \{\mathbf{0}\}$.
3. $I(p) = \overline{I}(p)$.

Proof. (1 $\Leftrightarrow$ 2). Let $V \in \mathcal{X}$. Then V is a semipositive motion if and only if $YV \in C(\mathcal{E})^+$ and $YV \neq \mathbf{0}$.

(2 $\Leftrightarrow$ 3). $YV \in Y(\mathcal{X}) \cap C(\mathcal{E})^+$ and $YV \neq \mathbf{0}$ if and only if $V \in I(p)$ and $V \notin \overline{I}(p)$. □

There is notion which is similar to, though weaker than, “bar equivalence”. Some tensegrities have semipositive motions, but no strictly positive ones (that is, $Y(\mathcal{X}) \cap \text{int } C(\mathcal{E})^+ = \emptyset$, see Lemma 3.3.2 on page 53).

Definition 2.1.12. If $G(p)$ has semipositive motions but no strictly positive motions, we call $G(p)$ *partially $\mathcal{X}$ -bar equivalent*.

In Corollary 3.6.2 on page 69 we show that if $G(p)$ is partially $\mathcal{X}$ -bar equivalent, then some portion of $G(p)$ is $\mathcal{X}$ -bar equivalent.

There’s one more concept we’ll need and that is the concept of stress. Remembering the prestressed concrete of Section 1.3 on page 4, we can see that this is a set of tensions on the cables and compressions on the struts that give a net zero force at each vertex. That works well for discrete tensegrities, but we’ll need a slightly more general definition in the future, so we’ll define stresses to be certain elements of $C^*(\mathcal{E})$, the topological dual of $C(\mathcal{E})$.

Definition 2.1.13. If X is a vector space, then its *topological dual*,² denoted X^* , is the space of all continuous linear functionals on X .

The nonnegative orthant of $C^*(\mathcal{E})$ will play a role later, so we'll give it the name $C^*(\mathcal{E})^+$.

Definition 2.1.14. If S is a set of elements in the vector space X , then the *annihilator* of S is the set

$$S^\perp = \{\mu \in X^* : \mu(s) = 0, \forall s \in S\}.$$

Definition 2.1.15. If S is a set of elements in the vector space X , then the *dual cone* of S is the set

$$S^* = \{\mu \in X^* : \mu(s) \geq 0, \forall s \in S\}.$$

Definition 2.1.16. A *stress* is an element $\mu \in C^*(\mathcal{E})$ such that $\mu \in Y(\mathcal{X})^\perp$. A *strictly positive stress* is a stress μ with $\mu(f) > 0$ for all nonzero $f \in C(\mathcal{E})^+$. A *semipositive stress* is a nonzero stress μ with $\mu(f) \geq 0$ for all $f \in C(\mathcal{E})^+$.

This choice of definition for “strictly positive” and “semipositive” is explained more fully by Proposition 3.3.17 on page 58.

Let's take a moment to see that this definition of stress matches our intuitive understanding in the finite case.

Proposition 2.1.17. *If $G(p)$ is a finite tensegrity with rigidity operator Y and $\mathcal{X} = \text{VF}(\mathcal{V})$, and $\mu \in C^*(\mathcal{E})$, then μ is a stress (in the sense of Definition 2.1.16) if and only if $Y^\top \mu = \mathbf{0}$.*

²This dual is “topological” in that it uses continuity, which requires a topology, in its definition. The *algebraic dual*, by contrast, is the set of all linear functionals on X .

Proof. Suppose, first, that μ is a stress. That is, $\mu(YV) = 0$ for all $V \in \mathcal{X} = \text{VF}(\mathcal{V})$. Then $(Y^\top \mu)V = 0$ for all V . But for $Y^\top \mu$ to have zero dot product with all $V \in \text{VF}(\mathcal{V})$, it must be the zero vector.

Conversely, suppose $Y^\top \mu = \mathbf{0}$. Then $\mu(YV) = (Y^\top \mu)V = \mathbf{0} \cdot V = 0$. $\square$

With our definitions in place, we can say concisely what Roth and Whiteley proved. The essence of their result is that having a strictly positive stress is the same as being bar equivalent (see Lemma 2.4.3 on page 35).

Definition 2.1.18. Here is a list of all the entities we have defined, in alphabetical order.

B, C, and S	Three pairwise-disjoint sets of two-element subsets of $\mathcal{V}$. Elements of B (<i>bars</i>) require their vertices to stay a fixed distance apart. Elements of S (<i>struts</i>) and C (<i>cables</i>) serve as lower and upper bounds on their vertices' distances respectively.
$\mathcal{C}$ and $\mathcal{S}$	$\mathcal{C} := C \cup B$ and $\mathcal{S} := S \cup B$.
$C(\mathcal{E})$	Continuous functions from $\mathcal{E}$ to $\mathbb{R}$. The <i>loads</i> , $Y(\mathcal{X})$, are a subspace of $C(\mathcal{E})$.
$C(\mathcal{E})^+$	$\{f \in C(\mathcal{E}) : f(e) \geq 0, \forall e \in \mathcal{E}\}$
$C(\mathcal{E})^-$	$\{f \in C(\mathcal{E}) : f(e) \leq 0, \forall e \in \mathcal{E}\}$
$C^*(\mathcal{E})$	Continuous linear functionals on $C(\mathcal{E})$, tensions and compressions on the edges.
$C^*(\mathcal{E})^+$	$\{\mu \in C^*(\mathcal{E}) : \mu(f) \geq 0, \forall f \in C(\mathcal{E})^+\}$.
$\mathcal{E}$	$\mathcal{S} \sqcup \mathcal{C}$. The edgeset. It receives its topology from $\mathbb{R}^n$ via $\mathcal{V}$, $\mathcal{V} \times \mathcal{V}$ and $(\mathcal{V} \times \mathcal{V})/\sim$ (see Section 3.2 on page 39) and is required to be compact.
$G(p)$	$\{\mathcal{V}; S, C, B; p\}$. A tensegrity.
$\overline{G}(p)$	$\{\mathcal{V}; \overline{B}; p\}$. The bar framework created by replacing all of $G(p)$'s edges with bars. That is, elements of $\overline{B}$ are bars and $\overline{B} = S \cup C \cup B$.
$I(p)$	$\{V \in \text{VF}(\mathcal{V}) : YV \in C(\mathcal{E})^+\}$. Motions of $G(p)$.
<i>continued on the next page</i>	

<i>continued from the previous page</i>	
$\bar{I}(p)$	$\{V \in \text{VF}(\mathcal{V}) : YV = \mathbf{0}\}$. Motions of $\bar{G}(p)$.
p	A 1-1 map from $\mathcal{V}$ to $\mathbb{R}^n$.
$T(p)$	The Euclidean motions. These can be induced from rigid motions of $\mathbb{R}^n$ and do not change the shape of the tensegrity (see Appendix A).
$\mathcal{V}$	The <i>vertices</i> of $G(p)$. A set with the topology induced from $\mathbb{R}^n$ via p . When $\mathcal{V}$ is a curve rather just a general set, we may call it γ .
$\text{VF}(\mathcal{V})$	Continuous vector fields on $\mathcal{V}$, the <i>variations</i> . Those variations which produce nonnegative loads are called <i>motions</i> .
$\mathcal{X}$	The <i>design variations</i> . A subspace of $\text{VF}(\mathcal{V})$ containing $T(p)$.
$Y: \text{VF}(\mathcal{V}) \rightarrow C(\mathcal{E})$	$\pm(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2))$ with the sign being $+$ for struts and $-$ for cables. The rigidity operator. Used to calculate the load for a given variation.

2.2 WORKING THE CROSSED-SQUARE EXAMPLE

[Figure 2.2.1](#) has our example again, this time with the edges numbered. The time has come to see what we can learn about it. We start by building the rigidity operator

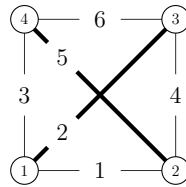


Figure 2.2.1: The crossed-square example with its edges numbered.

and choosing a stress:

$$Y = \begin{bmatrix} 1 & 0 & -1 & 0 & 0 & 0 & 0 & 0 \\ -1 & -1 & 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 1 & 0 & -1 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 0 & -1 & 1 \\ 0 & 0 & 0 & 0 & -1 & 0 & 1 & 0 \end{bmatrix} \quad \text{and} \quad \mu = \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \\ 1 \\ 1 \end{bmatrix}$$

Then $Y^\top \mu = \mathbf{0}$, so (by Proposition 2.1.17 on page 22) μ is a strictly positive stress.

By Lemma 2.4.3 on page 35, then, $G(p)$ is bar equivalent. It will be infinitesimally rigid or not as $\overline{G}(p)$ is.

Even though our main concern here is with bar equivalence, this example is simple enough that we can calculate whether $\overline{G}(p)$ is infinitesimally rigid or not. Since the primary value of bar equivalence is tied to knowing whether the bar framework is infinitesimally rigid, it seems worth the time to calculate it for this example. In Section 4.8 on page 108 we dig even deeper, showing that for this example, the image of Y is a hyperplane in $C(\mathcal{E})$.

Suppose that V is a motion of the bar framework $\overline{G}(p)$, that is, $V \in \overline{I}(p)$. Then $YV = \mathbf{0}$. But we have

$$YV = \begin{bmatrix} 1 & 0 & -1 & 0 & 0 & 0 & 0 & 0 \\ -1 & -1 & 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 1 & 0 & -1 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 0 & -1 & 1 \\ 0 & 0 & 0 & 0 & -1 & 0 & 1 & 0 \end{bmatrix} \begin{bmatrix} v_{x,1} \\ v_{y,1} \\ v_{x,2} \\ v_{y,2} \\ v_{x,3} \\ v_{y,3} \\ v_{x,4} \\ v_{y,4} \end{bmatrix} = \begin{bmatrix} v_{x,1} - v_{x,2} \\ -v_{x,1} - v_{y,1} + v_{x,3} + v_{y,3} \\ v_{y,1} - v_{y,4} \\ v_{y,2} - v_{y,3} \\ v_{x,2} - v_{y,2} - v_{x,4} + v_{y,4} \\ -v_{x,3} + v_{y,4} \end{bmatrix}$$

A little calculation gives us that any motion of the bar framework must look like

$$V = \begin{bmatrix} v_{x,1} & v_{y,1} & v_{x,1} & v_{y,2} & (v_{x,1} + v_{y,1} - v_{y,2}) & v_{y,2} & (v_{x,1} + v_{y,1} - v_{y,2}) & v_{y,1} \end{bmatrix}^\top$$

We'd like to know whether this variation is a Euclidean motion or whether it changes the shape of the bar framework. We can remove any translation from V by setting the sum of the horizontal components to zero and the sum of the vertical components to zero. That is we require that

$$v_{x,1} + v_{x,2} + v_{x,3} + v_{x,4} = 0 \quad \text{and} \quad v_{y,1} + v_{y,2} + v_{y,3} + v_{y,4} = 0.$$

If we do that, we get

$$V = \begin{bmatrix} v_{x,1} & -v_{x,1} & v_{x,1} & v_{x,1} & -v_{x,1} & v_{x,1} & -v_{x,1} & -v_{x,1} \end{bmatrix}^\top,$$

which is simply a rotation around the center of the figure (see [Figure 2.2.2](#)). So every

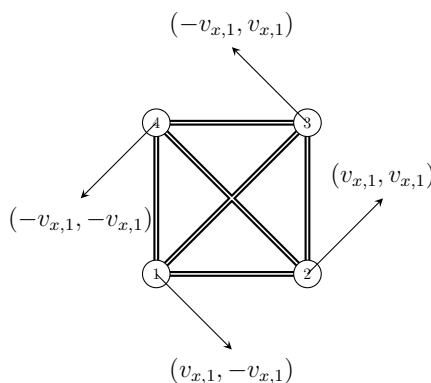


Figure 2.2.2: Every motion of the crossed-square bar framework is a Euclidean motion.

element of $\bar{I}(p)$ is an element of $T(p)$. That makes $\bar{G}(p)$ infinitesimally rigid and by Theorem 2.4.4 on page 35, $G(p)$ is infinitesimally rigid.

2.3 THE MECHANICS

In their proof, Roth and Whiteley make use of two theorems of Rockafellar, but he uses some terms we haven't seen before, so we'll mention them first.

Rockafellar uses the term “relative interior”. The idea is useful in a vector space when dealing with objects that would fit within proper subspaces of the vector space. For example, when looking at a line segment in $\mathbb{R}^3$, it is clear that there are two endpoints and many “non-end” points. However, a line segment in $\mathbb{R}^3$ has no 3-dimensional interior. That is where “relative interior” comes in handy. Those “non-end” points form the relative interior of the line segment (in the case of a single point, the point itself is the relative interior).

Let’s define that more rigorously. Our first two definitions differ in the bounds on λ .

Definition 2.3.1 (Rockafellar (1970, p. 10)). A subset C of $\mathbb{R}^n$ is called a *convex set* if $(1 - \lambda)x + \lambda y \in C$ whenever $x, y \in C$ and $\lambda \in (0, 1)$.

Definition 2.3.2 (Rockafellar (1970, p. 3)). A subset M of $\mathbb{R}^n$ is called an *affine set* if $(1 - \lambda)x + \lambda y \in M$ whenever $x, y \in M$ and $\lambda \in \mathbb{R}$.

Definition 2.3.3 (Rockafellar (1970, p. 6)). The *affine hull* of a set $S \subset \mathbb{R}^n$ is the intersection of the collection of affine sets M such that $M \supset S$.

Definition 2.3.4 (Rockafellar (1970, p. 44)). The *relative interior* of a convex set $C \subset \mathbb{R}^n$ is the set of points x in the affine hull of C for which there exists an $\varepsilon > 0$ such that $y \in C$ whenever y is in the affine hull of C and $d(x, y) \leq \varepsilon$.

For the sake of understanding, let’s apply these definitions to the line segment S running from $(0, 0, 0)$ to $(0, 0, 1)$ in $\mathbb{R}^3$.

Proposition 2.3.5. *The relative interior of the line segment $\{(0, 0, z) : z \in [0, 1]\}$ is the open set $\{(0, 0, z) : z \in (0, 1)\}$.*

Proof. The affine sets that contain S are the z -axis, any plane containing the z -axis and $\mathbb{R}^3$. The affine hull of S , then, is the intersection of those set, or simply the z -axis. Then, for any point $(0, 0, z)$, one of three things is true:

1. $z > 1$ or $z < 0$. In this case $(0, 0, z)$ itself is outside S , so no epsilon neighborhood of it in the affine hull could be completely contained within S .
2. $z = 0$ or $z = 1$. Here $(0, 0, z)$ is inside S , but for no positive ε is the neighborhood completely inside S , so these points are also not in the relative interior.
3. $z > 0$ and $z < 1$. In this case, setting $\varepsilon = \min\{1 - z, z - 0\}$ gives us a neighborhood in the affine hull that is also in S . These are the points of the relative interior. $\square$

While we're on the topic, let's support that offhand comment about the relative interior of a point.

Proposition 2.3.6. *The relative interior of a point is the point itself.*

Proof. Let $x \in \mathbb{R}^n$. Now the smallest (in terms of set inclusion) affine set containing x (and thus the affine hull of x) is x itself. But then, if y is in the affine hull, y must be x , so $d(x, y) = 0 \leq \varepsilon$ for any $\varepsilon > 0$. $\square$

Let's get back to definitions.

Definition 2.3.7 (Rockafellar (1970, p. 170)). A *polyhedral set* is one that is the intersection of a finite number of closed half-spaces.

Definition 2.3.8 (Rockafellar (1970, p. 162)). A *face* of a convex set C is a convex subset C' of C such that every (closed) line segment in C with a relative interior point in C' has both endpoints in C' .

Proposition 2.3.9. *For a cube in $\mathbb{R}^3$, the faces are: the 6 squares that form the boundary, the 12 edges of the squares and the 8 corners as well as the entire cube and the empty set.*

Proof. Clearly the empty set is trivially a face, as it contains no line segments. Similarly, the entire cube is a face as any closed line segment in the cube has both ends in the cube.

The 8 corners are faces in much the same fashion. They can hold endpoints of line segments, but the only way they can hold relative interior points of line segments is if those line segments are the zero-length ones which are the corners themselves, and in that case both endpoints are certainly in. No other single point can be a face, though, as any other point can be a relative interior point of some line segment that passes through it.

The 12 edges are faces, as any line segment that has a relative interior point in an edge must lie entirely in that edge. No other segment in the cube can be a face since it will be crossed by some other segment that it does not contain.

Finally, any segment with a relative interior point in one of the boundary squares of the cube must lie entirely in that square, but any other plane passing through the cube can be intersected transversely with some other segment in the cube, so those squares are the only 2-dimensional faces. $\square$

In considering unbounded convex sets, Rockafellar introduces a type of “point at infinity”. An unbounded convex set (not necessarily conic), he notes, must contain some entire half-line or it wouldn’t be an unbounded set. However, that half-line might point in any direction. So the directions in which the half-lines lie can be thought of as “horizon points” or points at infinity. Formally,

Definition 2.3.10 (Rockafellar (1970, p. 60)). A *direction* of $\mathbb{R}^n$ is an equivalence class of the collection of all closed half-lines of $\mathbb{R}^n$ under the equivalence relation “half-line L_1 is a translate of half-line L_2 .”

Definition 2.3.11 (Rockafellar (1970, p. 12)). The *convex hull* of a set $S \subset \mathbb{R}^n$ is the intersection of all convex sets containing S .

Definition 2.3.12 (Rockafellar (1970, p. 170)). A *finitely generated convex set* is the convex hull of a finite set of points and directions.

As an example, the convex hull in $\mathbb{R}^2$ of the two points $(0,0)$ and $(1,0)$ and the direction of the vector $(1,1)$ is the area shown in Figure 2.3.1.

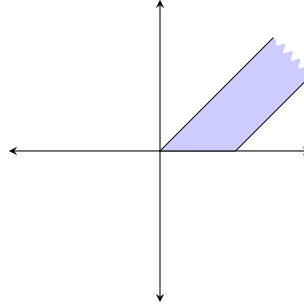


Figure 2.3.1: The convex set generated by the points $(0,0)$ and $(1,0)$ and the direction of the vector $(1,1)$.

Now we are ready for the theorems. The first is a separation theorem and the second gives an equivalence between attributes of a convex set.

Theorem 2.3.13. *Let X be a closed convex cone in $\mathbb{R}^n$ and $x_0 \in \mathbb{R}^n \setminus X$. Then, since the relative interiors of X and $\{x_0\}$ are disjoint, there exists a hyperplane separating X and $\{x_0\}$ properly.*

Proof. See Rockafellar (1970, Theorem 11.3, p. 97). □

Theorem 2.3.14. *The following properties of a convex set C are equivalent:*

- (a) C is polyhedral;
- (b) C is closed and has only finitely many faces;
- (c) C is finitely generated.

Proof. See Rockafellar (1970, Theorem 19.1, p. 171). □

We also have two more objects that will be of value to us. Suppose we have some set X of vectors in $\mathbb{R}^n$. It seems natural to consider the set of vectors that are normal

to all of our vectors. This is the finite-dimensional analogue of the annihilator, so we use the same notation:

$$X^\perp = \{v \in \mathbb{R}^n : v \cdot x = 0 \text{ for all } x \in X\}$$

Similarly, the set of vectors that “point in the same general direction” as ours is the analog of the dual cone:

$$X^* = \{v \in \mathbb{R}^n : v \cdot x \geq 0 \text{ for all } x \in X\}.$$

Lemma 2.3.15. *For any set $X \subset \mathbb{R}^n$, X^* is a closed convex cone.*

Proof. If $x_1, x_2 \in X$, then, for any $x \in X$,

$$((1 - \lambda)x_1 + \lambda x_2) \cdot x = (1 - \lambda)x_1 \cdot x + \lambda x_2 \cdot x \geq 0$$

whenever $\lambda \in [0, 1]$. Also, $\alpha x_1 \cdot x \geq 0$ whenever $\alpha > 0$, so X^* is a convex cone. If y lies outside X^* , then by the definition of X^* there must be some $x_- \in X$ such that $y \cdot x_- < 0$. Let $z \in \mathbb{R}^n$ with $\|y - z\| < \frac{-y \cdot x_-}{\|x_-\|}$. Then we have

$$\begin{aligned} z \cdot x_- &= (z - y + y) \cdot x_- = (z - y) \cdot x_- + y \cdot x_- \\ &\leq (z - y) \cdot (z - y) \frac{\|x_-\|}{\|z - y\|} + y \cdot x_- = \|y - z\| \|x_-\| + y \cdot x_- \\ &< -y \cdot x_- + y \cdot x_- = 0. \end{aligned}$$

That is, z is also outside X^* . So X^* is closed. □

Here’s a corollary to Theorem 2.3.13 that will prove useful.

Corollary 2.3.16. *If X is a closed convex cone in $\mathbb{R}^n$ and $x_0 \in \mathbb{R}^n \setminus X$, then there exists $\mu \in X^*$ with $\mu \cdot x_0 < 0$.*

Proof. By Theorem 2.3.13 there exists a hyperplane, h , with x_0 lying strictly on one side of it and X in the half-space on the other side. The normal to h pointing into the half-space with X has nonnegative dot product with all the elements of X and hence is in X^* , but it has negative dot product with x_0 . $\square$

Here are a few more attributes of X^* :

Proposition 2.3.17.

- (i) for $Z \subset \mathbb{R}^n$ and X a closed convex cone in $\mathbb{R}^n$, $Z \subset X$ if and only if $X^* \subset Z^*$.
- (ii) for $Z \subset \mathbb{R}^n$, $Z^{**} = (Z^*)^*$ is the smallest closed convex cone in $\mathbb{R}^n$ that contains Z .
- (iii) if $Z = \{z_1, \dots, z_k\}$ is a finite subset of $\mathbb{R}^n$, then the convex cone

$$\left\{ \sum_{i=1}^k \lambda_i z_i : \lambda_i \geq 0 \text{ for } 1 \leq i \leq k \right\}$$

generated by Z is closed in $\mathbb{R}^n$ (and therefore equals Z^{**}).

Proof.

- (i) Suppose that X is some closed convex cone in $\mathbb{R}^n$ and we have some $Z \subset \mathbb{R}^n$. If $Z \subset X$, then certainly any vector that has a positive dot product with everything in X must have a positive dot product with everything in Z , hence $X^* \subset Z^*$. What about the other way around? If there is some $z \in Z, z \notin X$, then by Corollary 2.3.16 on the preceding page, there is some $\mu \in X^*$ such that $\mu \cdot z < 0$. But that means that $\mu \notin Z^*$.
- (ii) Z^{**} is a closed convex cone by Lemma 2.3.15. Furthermore, if $z \in Z$, then z has a positive dot product with everything in Z^* , so $Z \subset Z^{**}$. Suppose that some closed convex cone $C \supset Z$ but that there were some $\hat{z} \in Z^{**}$ and $\hat{z} \notin C$. Then, by Corollary 2.3.16, there must be an element μ of C^* that separates

$\hat{z}$ and C . Now by Item (i), $C^* \subset Z^*$, so $\mu \in Z^*$ and yet $\mu \cdot \hat{z} < 0$. That contradicts the definition of Z^{**} . Hence Z^{**} must be the smallest closed convex cone containing Z .

(iii) Finally, if Z is a finite subset of $\mathbb{R}^n$, then the convex cone

$$Z_c = \left\{ \sum_{i=1}^k \lambda_i z_i : \lambda_i \geq 0 \text{ for } 1 \leq i \leq k \right\}$$

is finitely generated (being generated by the origin and the directions of the z_i) and hence, by Theorem 2.3.14 on page 30, it is closed. Now any positive linear combination of the z_i must have positive dot product with any vector in Z^* , so $Z_c \subset Z^{**}$. Since it is a closed convex cone it must, by Item (ii), equal Z^{**} . $\square$

One lemma remains that we will want in the next section.

Lemma 2.3.18. *Let $X = \{x_1, \dots, x_k\} \in \mathbb{R}^n$. Then $X^* = X^\perp$ implies that X^{**} is a subspace.*

Proof. Let $\alpha \in \mathbb{R}$ and $v_1, v_2 \in X^{**}$. We want to show that $v_1 + v_2 \in X^{**}$ and that $\alpha v_1 \in X^{**}$.

$$\begin{aligned} v_1, v_2 \in X^{**} &\Rightarrow v_1 \cdot x^* \geq 0 \text{ and } v_2 \cdot x^* \geq 0, & \forall x^* \in X^* \\ &\Rightarrow (v_1 + v_2) \cdot x^* = v_1 \cdot x^* + v_2 \cdot x^* \geq 0, & \forall x^* \in X^* \\ &\Rightarrow (v_1 + v_2) \in X^{**} \end{aligned}$$

Now if $\alpha \geq 0$, then we easily have that $(\alpha v_1) \cdot x^* = \alpha v_1 \cdot x^* \geq 0$ for all $x^* \in X^*$. But if $\alpha < 0$, we need to be a little more clever. In that case, we have $\alpha v_1 \cdot x^* = |\alpha| v_1 \cdot -x^*$, so we are fine if $x^* \in X^* \Rightarrow -x^* \in X^*$. But

$$\begin{aligned} x^* \in X^* &\Rightarrow x^* \cdot x = 0, \forall x \in X, \text{ because } X^* = X^\perp \text{ by hypothesis} \\ &\Rightarrow -x^* \cdot x = 0, \forall x \in X \\ &\Rightarrow -x^* \in X^* \end{aligned}$$

and we have that X^{**} is a subspace. $\square$

2.4 KEY LEMMA AND THEOREM

We now reach the key lemma (Roth and Whiteley, 1981, Lemma 5.1, p. 426). This is the engine that Roth & Whiteley use to establish bar equivalence for a tensegrity so that they can analyze the infinitesimal rigidity of the bar framework instead of the tensegrity. It is the predecessor to our main theorem. The proof we give here is essentially that of Roth & Whiteley.

Lemma 2.4.1. *Suppose $X = \{x_1, \dots, x_k\} \subset \mathbb{R}^n$. Then $X^* = X^\perp$ if and only if there exist positive scalars $\lambda_1, \dots, \lambda_k$ with $\sum_{i=1}^k \lambda_i x_i = \mathbf{0}$.*

Proof. By Item (iii) of Proposition 2.3.17, $X^{**} = \left\{ \sum_{i=1}^k \lambda_i x_i : \lambda_i \geq 0 \text{ for } 1 \leq i \leq k \right\}$.

Suppose that $X^* = X^\perp$. Then, by Lemma 2.3.18, X^{**} is a subspace. Since every $x_j \in X$ is also in X^{**} we can write

$$-x_j = \sum_{i=1}^k \lambda_i x_i$$

for some set of $\lambda_i \geq 0$. But then we have

$$\lambda_1 x_1 + \dots + (1 + \lambda_j) x_j + \dots + \lambda_k x_k = \mathbf{0}.$$

Adding up k such expressions gives a linear dependency of $x_1, \dots, x_k$ with all positive coefficients.

Conversely, if there is some set $\lambda_i > 0$ for which $\sum_{i=1}^k \lambda_i x_i = \mathbf{0}$ and if $x^* \in X^*$, then

$$0 = x^* \cdot \mathbf{0} = x^* \cdot \sum_{i=1}^k \lambda_i x_i = \sum_{i=1}^k \lambda_i (x^* \cdot x_i).$$

Since $x^* \cdot x_i \geq 0$ and $\lambda_i > 0$ for all i , we must have $x^* \cdot x_i = 0$ for all i , so $x^* \in X^\perp$. $\square$

Roth & Whiteley's lemma is very much in the world of finite-dimensional Euclidean spaces. We need to reformulate it in preparation for the move to more general spaces.

Lemma 2.4.2. *Let Y be an $m \times n$ matrix with elements in $\mathbb{R}$. Then exactly one of the following is true:*

(I) *There exists $V \in \mathbb{R}^n$ such that $YV \geq \mathbf{0}$.*

(II) *There exists $\mu \in \mathbb{R}^m$ with $\mu > \mathbf{0}$ such that $Y^\top \mu = \mathbf{0}$.*

Proof. Let $X = \{x_i\}$ be the set of row vectors of Y . Then those V for which $YV = \mathbf{0}$ make up $X^\perp$ and those V for which YV has no negative coordinates make up X^* .

Case (I) is false if and only if $X^* = X^\perp$. By Lemma 2.4.1, $X^* = X^\perp$ if and only if there exist positive scalars $\lambda_1, \dots, \lambda_m$ such that $\sum_{i=1}^m \lambda_i x_i = \mathbf{0}$. And by defining $\mu = \begin{bmatrix} \lambda_1 \\ \vdots \\ \lambda_m \end{bmatrix}$ we have that the λ_i 's exist if and only if Case (II) is true. $\square$

Lemma 2.4.2 resembles the “Theorem of the Alternative” we talk about in Section 3.1, but it doesn't provide a lot of geometric insight. Let's recast it yet one more time.

Lemma 2.4.3. *Let $G(p)$ be a (finite) tensegrity framework with rigidity matrix Y . Then $G(p)$ is bar equivalent if and only if $G(p)$ has a strictly positive stress.*

Proof. Let $G(p)$ be bar equivalent. Then there is no V for which $YV \geq \mathbf{0}$, so we are not in Case (I) of Lemma 2.4.2. Hence we must be in Case (II), i.e., $G(p)$ has a strictly positive stress.

Conversely, let $G(p)$ have a strictly positive stress. We are again in Case (II), so there can be no V like that described in Case (I), and $G(p)$ is bar equivalent. $\square$

Finally, as an application of Lemma 2.4.3, we prove Roth and Whiteley's theorem (1981, Theorem 5.2, p. 427).

Theorem 2.4.4. *Suppose $G(p)$ is a (finite) tensegrity framework in $\mathbb{R}^n$ and $\overline{G}(p)$ the associated bar framework. Then $G(p)$ is infinitesimally rigid if and only if $\overline{G}(p)$ is infinitesimally rigid and there exists a strictly positive stress of $G(p)$.*

Proof. ($\Rightarrow$). Suppose that $G(p)$ is infinitesimally rigid. Then $I(p) = \bar{I}(p) = T(p)$ by Definition 2.1.9 on page 20. So $\bar{G}(p)$ is infinitesimally rigid (again by Definition 2.1.9) and $G(p)$ is bar equivalent by Proposition 2.1.10. By Lemma 2.4.3, $G(p)$ has a strictly positive stress.

($\Leftarrow$). Conversely, suppose that $\bar{G}(p)$ is infinitesimally rigid and that $G(p)$ has a strictly positive stress. By Lemma 2.4.3, $G(p)$ is bar equivalent. Infinitesimal rigidity of $\bar{G}(p)$ gives us $\bar{I}(p) = T(p)$ and bar equivalence of $G(p)$ gives us $I(p) = \bar{I}(p)$, so $I(p) = T(p)$ and thus $G(p)$ is infinitesimally rigid. $\square$

CHAPTER 3

MAIN THEOREMS

3.1 UNDER THE HOOD: THEOREMS OF THE ALTERNATIVE

Having seen what Roth & Whiteley accomplished, we'd like to get a better grasp of what makes their theorem work, so that we can extend it to continuous tensegrities. Let's take another look at Lemma 2.4.2 on page 35. The statement that $Y^\top \mu = \mathbf{0}$ implies that $(Y^\top \mu)V = 0$ for all $V \in \text{Vf}(\mathcal{V})$. That's the same as saying that $\mu^\top YV = 0$ for all V or that $\mu \in (\text{im } Y)^\perp$.

With that understanding, we can see Lemma 2.4.2 as saying, “either $\text{im } Y$ intersects the nonnegative orthant somewhere other than the origin, or else $(\text{im } Y)^\perp$ intersects the interior of the nonnegative orthant, but not both” (see Figure 3.1.1).

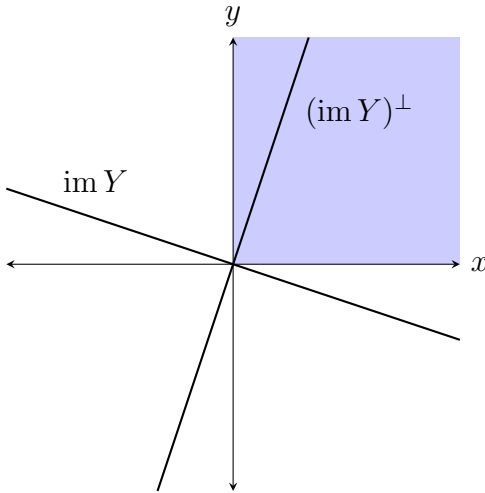


Figure 3.1.1: A Theorem of the Alternative in $\mathbb{R}^2$.

It turns out that this is one of many “Theorems of the Alternative” (also sometimes called “Theorems on the Alternative”), in which exactly one of two alternative systems

of inequalities or equalities is shown to have a solution. In his *Nonlinear Programming*, Olvi Mangasarian has a wonderful chapter where he catalogs no fewer than 11 of these theorems (Mangasarian, 1969, ch. 2). Mangasarian indicates that the one we have in hand is the one from Stiemke's 1915 paper (Stiemke, 1915).¹ It is very similar to a theorem he credits to Paul Gordan (1873). Here are Mangasarian's statements of the two, using our notation.

Theorem 3.1.1 (Stiemke 1915). *For a given matrix Y , either*

$$(I) \ YV \not\geq \mathbf{0} \text{ has a solution } V$$

or

$$(II) \ Y^\top \mu = \mathbf{0}, \mu > \mathbf{0} \text{ has a solution } \mu$$

but never both.

Theorem 3.1.2 (Gordan 1873). *For a given matrix Y , either*

$$(I) \ YV > \mathbf{0} \text{ has a solution } V$$

or

$$(II) \ Y^\top \mu = \mathbf{0}, \mu \not\geq \mathbf{0} \text{ has a solution } \mu$$

but never both.

The difference between those two is in which case is true when both $\text{im } Y$ and $(\text{im } Y)^\perp$ lie on the boundar of the nonnegative orthant. That difference will appear again later in our main theorems. Before we get there, though, we need to go back through the various entities listed in Definition 2.1.18 on pages 23–24 and see how they change during the move to the continuous world.

¹The papers of Stiemke, Gordan and Motzkin (see Appendix B) are all in German, so, due to time constraints, I have not verified Mangasarian's analysis of them.

3.2 TURNING SETS INTO SPACES: $\mathcal{V}$ AND $\mathcal{E}$

3.2.1 NEW TOPOLOGIES VIA METRICS

We will need $\mathcal{E}$ to be a compact metric space, so we set out to provide a metric on the set. Now $\mathcal{E}$ is, by definition, a disjoint union of two sets. Each of those sets is a subset of the quotient space $(\mathcal{V} \times \mathcal{V})/\sim$ formed from the direct product $\mathcal{V} \times \mathcal{V}$ by the identification $(v_1, v_2) \sim (v_2, v_1)$. We'll work our way through these spaces in creating the desired metric.

Definition 3.2.1. The space $\mathcal{V}$ is the set of vertices, the metric induced by its embedding into $\mathbb{R}^n$:

$$d_{\mathcal{V}}(v_1, v_2) := \|p(v_1) - p(v_2)\|$$

and the topology defined by $d_{\mathcal{V}}$.²

Now, a metric has three attributes. It must be positive definite and symmetric, and it must satisfy the triangle inequality. Of course $d_{\mathcal{V}}$ has those attributes naturally, but we'll need to check them for the metrics yet to come.

We give $\mathcal{V} \times \mathcal{V}$ the metric

$$d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_3, v_4)) := \max\{d_{\mathcal{V}}(v_1, v_3), d_{\mathcal{V}}(v_2, v_4)\}.$$

This has the decided advantage that the basic open sets are the basic open sets of the product topology. For example, if we take the basic open set of size ε around (v_1, v_2) , then the points inside it are the edges (v_3, v_4) where $d_{\mathcal{V}}(v_1, v_3) < \varepsilon$ and $d_{\mathcal{V}}(v_2, v_4) < \varepsilon$.

Proposition 3.2.2. $d_{\mathcal{V} \times \mathcal{V}}$ is a metric.

Proof. Now $d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_3, v_4))$ is clearly nonnegative and is zero only when $v_1 = v_3$ and $v_2 = v_4$. Its symmetry comes from the nature of max. So we only need to check the triangle inequality.

²That is, the topology whose basic open sets are the open balls $d_{\mathcal{V}}(v_1, v_2) < \varepsilon$ for all $v_1, v_2 \in V$, $\varepsilon > 0$. See, for example [Munkres \(2000, p. 119\)](#).

If $(v_1, v_2), (v_3, v_4), (v_5, v_6) \in \mathcal{V} \times \mathcal{V}$, then

$$d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_5, v_6)) = \max\{d_{\mathcal{V}}(v_1, v_5), d_{\mathcal{V}}(v_2, v_6)\}.$$

Now, if $d_{\mathcal{V}}(v_1, v_5) \geq d_{\mathcal{V}}(v_2, v_6)$, we have

$$\begin{aligned} d_{\mathcal{V}}(v_2, v_6) &\leq d_{\mathcal{V}}(v_1, v_5) \leq d_{\mathcal{V}}(v_1, v_3) + d_{\mathcal{V}}(v_3, v_5) \\ &\leq \max\{d_{\mathcal{V}}(v_1, v_3), d_{\mathcal{V}}(v_2, v_4)\} + \max\{d_{\mathcal{V}}(v_3, v_5), d_{\mathcal{V}}(v_4, v_6)\}. \end{aligned}$$

Otherwise, we have

$$\begin{aligned} d_{\mathcal{V}}(v_1, v_5) &\leq d_{\mathcal{V}}(v_2, v_6) \leq d_{\mathcal{V}}(v_2, v_4) + d_{\mathcal{V}}(v_4, v_6) \\ &\leq \max\{d_{\mathcal{V}}(v_1, v_3), d_{\mathcal{V}}(v_2, v_4)\} + \max\{d_{\mathcal{V}}(v_3, v_5), d_{\mathcal{V}}(v_4, v_6)\}. \end{aligned}$$

So, either way, we have

$$d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_5, v_6)) \leq d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_3, v_4)) + d_{\mathcal{V} \times \mathcal{V}}((v_3, v_4), (v_5, v_6)).$$

Thus $d_{\mathcal{V} \times \mathcal{V}}$ is a metric. □

Now we come to $(\mathcal{V} \times \mathcal{V})_{/\sim}$.

Proposition 3.2.3. *The function*

$$d_{(\mathcal{V} \times \mathcal{V})_{/\sim}}(\{v_1, v_2\}, \{v_3, v_4\}) := \min_{i \neq j \in \{3,4\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_i, v_j))\}$$

is a metric for $(\mathcal{V} \times \mathcal{V})_{/\sim}$.

Proof. Let $\{v_1, v_2\}$, $\{v_3, v_4\}$, and $\{v_5, v_6\}$ be in $(\mathcal{V} \times \mathcal{V})_{/\sim}$.

- (a) $d_{(\mathcal{V} \times \mathcal{V})_{/\sim}}(\{v_1, v_2\}, \{v_3, v_4\})$ is the smaller of two nonnegative numbers and hence nonnegative. Furthermore, as $d_{\mathcal{V} \times \mathcal{V}}$ is a metric, $d_{(\mathcal{V} \times \mathcal{V})_{/\sim}}$ will only return 0 when either $(v_1, v_2) = (v_3, v_4)$ or $(v_1, v_2) = (v_4, v_3)$. That is, only when the two are in the same equivalence class in $(\mathcal{V} \times \mathcal{V})_{/\sim}$.

(b) We need symmetry of the metric:

$$\begin{aligned}
d_{(\mathcal{V} \times \mathcal{V})/\sim}(\{v_1, v_2\}, \{v_3, v_4\}) &= \min_{i \neq j \in \{3,4\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_i, v_j))\} \\
&= \min_{i \neq j \in \{3,4\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_i, v_j), (v_1, v_2))\} \\
&= \min_{i \neq j \in \{1,2\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_3, v_4), (v_i, v_j))\} \\
&= d_{(\mathcal{V} \times \mathcal{V})/\sim}(\{v_3, v_4\}, \{v_1, v_2\})
\end{aligned}$$

(c) Finally, we need to check the triangle inequality:

$$\begin{aligned}
d_{(\mathcal{V} \times \mathcal{V})/\sim}(\{v_1, v_2\}, \{v_5, v_6\}) &= \min_{i \neq j \in \{5,6\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_i, v_j))\} \\
&\leq \min_{\substack{s \neq t \in \{3,4\} \\ i \neq j \in \{5,6\}}} \{d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_s, v_t)) + d_{\mathcal{V} \times \mathcal{V}}((v_s, v_t), (v_i, v_j))\} \\
&= \min_{s \neq t \in \{3,4\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_s, v_t))\} + \min_{i \neq j \in \{5,6\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_3, v_4), (v_i, v_j))\} \\
&= d_{(\mathcal{V} \times \mathcal{V})/\sim}(\{v_1, v_2\}, \{v_3, v_4\}) + d_{(\mathcal{V} \times \mathcal{V})/\sim}(\{v_3, v_4\}, \{v_5, v_6\}).
\end{aligned}$$

So $d_{(\mathcal{V} \times \mathcal{V})/\sim}$ is a metric for $(\mathcal{V} \times \mathcal{V})/\sim$ and also for $\mathcal{S}$ and $\mathcal{C}$ as subsets of $(\mathcal{V} \times \mathcal{V})/\sim$. □

We've nearly achieved our goal, but we have a decision to make. How far apart should a strut be from a cable? Both $\mathcal{S}$ and $\mathcal{C}$ are subsets of $(\mathcal{V} \times \mathcal{V})/\sim$, so we could simply define the distance from strut s to cable c to be $d_{(\mathcal{V} \times \mathcal{V})/\sim}(s, c)$.

At first glance this seems a wise idea. After all, suppose we have a tensegrity in which the struts or the cables alone do not form a compact set, but the two together do (such a construction may be seen in [Figure 3.2.1](#) on the following page).³ If we

³Since $\mathcal{C}$ is Hausdorff, we know that if $\mathcal{S}$ and $\mathcal{C}$ were compact, they would be closed (see, for example, [Munkres \(2000, p. 165\)](#)). But the sequence of struts at angles $\pi - \frac{1}{2^n}$ and the sequence of cables at angles $-\frac{1}{2^n}$ show that we would need a strut at π and a cable at 0 for closure.

On the other hand, $\mathcal{V}$ is compact as a $p(\mathcal{V})$ is a closed, bounded subset of $\mathbb{R}^2$, so $\mathcal{V} \times \mathcal{V}$ is compact. Now $\mathcal{S} \cup \mathcal{C}$ is closed, so if $\pi: \mathcal{V} \times \mathcal{V} \rightarrow (\mathcal{V} \times \mathcal{V})/\sim$ is the quotient map, then $\pi^{-1}(\mathcal{S} \cup \mathcal{C})$ is a closed subset of $\mathcal{V} \times \mathcal{V}$ and hence compact. But then $\mathcal{S} \cup \mathcal{C}$ is the continuous image of a compact set and thus also compact ([Munkres, 2000, p. 165–167](#)).

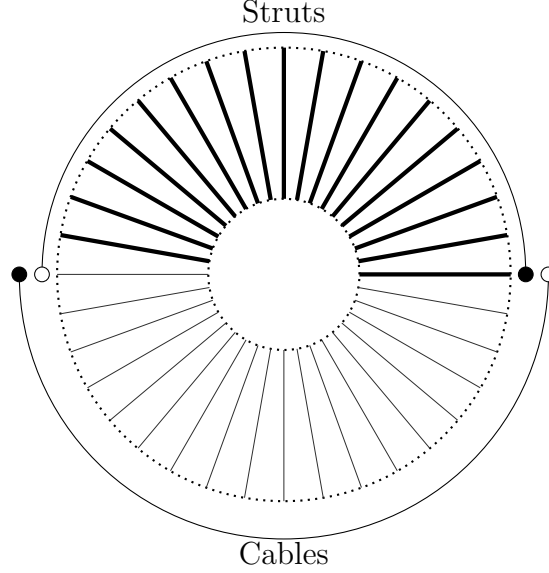


Figure 3.2.1: A tensegrity with noncompact $\mathcal{S}$ and $\mathcal{C}$ but compact $\mathcal{S} \cup \mathcal{C}$.

make the distance between struts and cables simply the $d_{(\mathcal{V} \times \mathcal{V})/\sim}$ distance, then this edgeset will be compact and our results here will apply to it.

However, if we do that and $\mathbf{B}$ is nonempty, we run into trouble. If b_s is a bar in $\mathcal{S}$ and b_c is the same bar in $\mathcal{C}$, then $d_{(\mathcal{V} \times \mathcal{V})/\sim}(b_s, b_c)$ would be zero. Putting the metric topology on $\mathcal{E}$ in this situation would mean that there would be no open sets which would separate b_s and b_c . $\mathcal{E}$ would no longer be Hausdorff. As we certainly need $\mathcal{E}$ to be Hausdorff, we'll look for other options (and we'll circle around later and talk for a bit about why excluding this example doesn't trouble us greatly).

Since having a strut and cable too close together was a problem, perhaps we should define them to be far apart. To do this, we could find out how far apart any two struts or cables are and then make all struts farther away from all cables than that. This is not difficult, but it seems clearer to use what [Tabor and Tabor \(2002\)](#) call an “extended metric” (they credit this to [Covitz and Nadler \(1970\)](#), who pass the credit on to [Luxemburg \(1958\)](#)).

An extended metric differs from a normal metric only in that it is allowed to take the value ∞ (to readers interested in the properties of metrics are recommended [Wikipedia \(2007b\)](#) and [Munkres \(2000, p. 119\)](#)).

Proposition 3.2.4. *The set $\mathcal{S} \sqcup \mathcal{C}$ is a metric space with (extended) metric $d_{\mathcal{E}}$ defined as:*

$$d_{\mathcal{E}}(e_1, e_2) = \begin{cases} d_{(\mathcal{V} \times \mathcal{V})/\sim}(e_1, e_2), & \text{if } e_1, e_2 \in \mathcal{S} \text{ or } e_1, e_2 \in \mathcal{C}, \text{ and} \\ \infty & \text{otherwise.} \end{cases}$$

Proof. Let e_1, e_2 and e_3 be in $\mathcal{S} \sqcup \mathcal{C}$. Then, if e_1 and e_2 are both struts or both cables, we get that $d_{\mathcal{E}}(e_1, e_2) \geq 0$ with equality only if $e_1 = e_2$ and that $d_{\mathcal{E}}(e_1, e_2) = d_{\mathcal{E}}(e_2, e_1)$, both from the nature of $d_{(\mathcal{V} \times \mathcal{V})/\sim}$. If they are different, then $d_{\mathcal{E}}(e_1, e_2) = d_{\mathcal{E}}(e_2, e_1) = \infty > 0$. So we are left with the triangle inequality to check.

Again, if all three edges are of the same type,

$$d_{\mathcal{E}}(e_1, e_2) + d_{\mathcal{E}}(e_2, e_3) \geq d_{\mathcal{E}}(e_1, e_3) \quad (3.1)$$

due to that being true for $d_{(\mathcal{V} \times \mathcal{V})/\sim}$. Otherwise, at least two differ in type. If e_1 differs from e_3 , then the right side of [Equation \(3.1\)](#) is ∞ , but one of the terms on the left must be as well (e_1 and e_3 are different, so e_2 must differ from one or the other). On the other hand, if e_1 and e_3 are of the same type, the right side is finite and both of the terms on the left are infinite. We see that the triangle inequality holds for $d_{\mathcal{E}}$.

So $d_{\mathcal{E}}$ is a metric on $\mathcal{S} \sqcup \mathcal{C}$. □

Definition 3.2.5. The *edgeset* $\mathcal{E}$ is the space formed by endowing the set $\mathcal{S} \sqcup \mathcal{C}$ with the metric $d_{\mathcal{E}}$ and the associated metric topology.

3.2.2 FAMILIAR TOPOLOGIES

Now that we have topologies on $\mathcal{V}$ and $\mathcal{E}$, let's see what we can learn about them.

Proposition 3.2.6. *The topology on $\mathcal{V}$, $\mathcal{T}_{\mathcal{V}}$ is same as the topology, $\mathcal{T}_{p(\mathcal{V})}$, which $\mathcal{V}$ would have if we put the subspace topology on $p(\mathcal{V})$ and then mapped all of the open sets back to $\mathcal{V}$ via p .*

Proof. We need to show that the basic open sets in each topology are open in the other topology.

Basic open sets in $\mathcal{T}_{\mathcal{V}}$ are induced by the intersection with $p(\mathcal{V})$ of open balls in $\mathbb{R}^n$ centered at points of $p(\mathcal{V})$. Basic open sets in $\mathcal{T}_{p(\mathcal{V})}$ are induced by the intersection with $p(\mathcal{V})$ of open balls in $\mathbb{R}^n$ centered anywhere, so $\mathcal{T}_{\mathcal{V}} \subset \mathcal{T}_{p(\mathcal{V})}$.

To show the converse, let B be a basic open set in $\mathcal{T}_{p(\mathcal{V})}$ containing the vertex v , and let $r \in (0, \infty)$ and $c \in \mathbb{R}^n$ be the radius and center, respectively, of an n -ball which, when intersected with $p(\mathcal{V})$ gives $p(B)$. Since $v \in B$, we know that $\|c - p(v)\| < r$, so the intersection with $p(\mathcal{V})$ of the open n -ball of radius $r - \|c - p(v)\|$ centered at $p(v)$ induces an open set in $\mathcal{T}_{\mathcal{V}}$ (by definition) which contains v (by design) and is contained in B because if $x \in \mathbb{R}^n$ with $\|p(v) - x\| < r - \|c - p(v)\|$ we have

$$\|c - x\| = \|c - p(v) + p(v) - x\| \leq \|c - p(v)\| + \|p(v) - x\| < r$$

(see [Figure 3.2.2](#)).

□

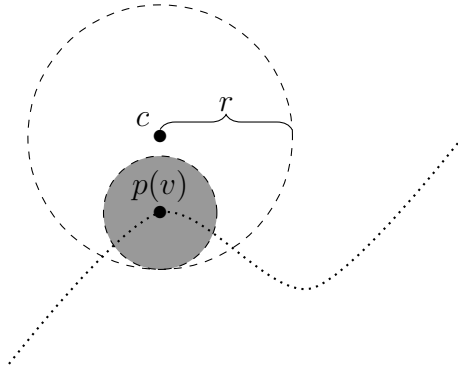


Figure 3.2.2: Open sets in $\mathcal{T}_{p(\mathcal{V})}$ are open in $\mathcal{T}_{\mathcal{V}}$.

So the topology in $\mathcal{V}$ is familiar. What about the one on $\mathcal{E}$? Let $\mathcal{T}_{\mathcal{S}}$ and $\mathcal{T}_{\mathcal{C}}$ be the topologies on $\mathcal{S}$ and $\mathcal{C}$ which come from seeing each of those spaces as a subspace

of a quotient space of $\mathcal{V} \times \mathcal{V}$. Then let

$$\mathcal{T}_{sqp} = \{U_{\mathcal{S}} \sqcup U_{\mathcal{C}} : U_{\mathcal{S}} \in \mathcal{T}_{\mathcal{S}}, U_{\mathcal{C}} \in \mathcal{T}_{\mathcal{C}}\}.$$

Proposition 3.2.7. $\mathcal{T}_{sqp}$ is a topology.

Proof. We need to show that $\mathcal{T}_{sqp}$ contains the empty set, the set $\mathcal{E}$, and that it is closed under arbitrary unions and finite intersections.

We note that the empty set is in $\mathcal{T}_{sqp}$ as it is the union of the empty set in $\mathcal{T}_{\mathcal{S}}$ and the empty set in $\mathcal{T}_{\mathcal{C}}$. Similarly, the set $\mathcal{E}$ is in $\mathcal{T}_{sqp}$ as it is the union of the sets $\mathcal{S} \in \mathcal{T}_{\mathcal{S}}$ and $\mathcal{C} \in \mathcal{T}_{\mathcal{C}}$. If $\{U_{\mathcal{S}_\alpha} \sqcup U_{\mathcal{C}_\alpha}\}$ is a collection of elements of $\mathcal{T}_{sqp}$, then

$$\bigcup \{U_{\mathcal{S}_\alpha} \sqcup U_{\mathcal{C}_\alpha}\} = \bigcup U_{\mathcal{S}_\alpha} \sqcup \bigcup U_{\mathcal{C}_\alpha}$$

is also in $\mathcal{T}_{sqp}$. Likewise, if $\{U_{\mathcal{S}_1} \sqcup U_{\mathcal{C}_1}, \dots, U_{\mathcal{S}_n} \sqcup U_{\mathcal{C}_n}\}$ is a finite collection of elements of $\mathcal{T}_{sqp}$, then

$$\bigcap \{U_{\mathcal{S}_i} \sqcup U_{\mathcal{C}_i}\} = \bigcap U_{\mathcal{S}_i} \sqcup \bigcap U_{\mathcal{C}_i}$$

is in $\mathcal{T}_{sqp}$ as well. So $\mathcal{T}_{sqp}$ is a topology for $\mathcal{E}$. □

Proposition 3.2.8. $\mathcal{T}_{\mathcal{E}}$ equals $\mathcal{T}_{sqp}$.

Proof. We note that the basic open sets of $\mathcal{T}_{sqp}$ consist of all edges in $\mathcal{E}$ whose ends fall within open sets in $\mathcal{V}$, U_1 and U_2 . We'll call such a set $U_1 \times U_2 / \sim$.

Let U_1 be an open set in $\mathcal{T}_{\mathcal{E}}$ containing edge $e_1 = \{v_{11}, v_{12}\}$. Without loss of generality, we can take $e_1 \in \mathcal{S}$. Then there is a basic open set, B_1 , centered at e_1 and of radius $0 < r_1 < \infty$ that is contained in U_1 . Since r_1 is finite, we know that the elements of B_1 are also in $\mathcal{S}$.

Let B_{11} and B_{12} be open balls in $\mathcal{V}$ centered at v_{11} and v_{12} , respectively, both of radius r_1 . Then, if $e_2 = \{v_{21}, v_{22}\} \in B_{11} \times B_{12} / \sim$, we have

$$\begin{aligned} d_{\mathcal{E}}(e_1, e_2) &= d_{(\mathcal{V} \times \mathcal{V}) / \sim}(e_1, e_2) \\ &= \min_{i \neq j \in \{1,2\}} \{d_{\mathcal{V} \times \mathcal{V}}((v_{11}, v_{12}), (v_{2i}, v_{2j}))\} \\ &\leq d_{\mathcal{V} \times \mathcal{V}}((v_{11}, v_{12}), (v_{21}, v_{22})) \\ &= \max\{d_{\mathcal{V}}(v_{11}, v_{21}), d_{\mathcal{V}}(v_{12}, v_{22})\} < r_1. \end{aligned}$$

so $e_2 \in B_1$. That is, the open set $B_{11} \times B_{12} / \sim \cap \mathcal{E}$ is contained in B_1 . So B_1 is open in $\mathcal{T}_{sqp}$.

Conversely, let U_{sqp} be an open set in $\mathcal{T}_{sqp}$ and let $e_1 = \{v_{11}, v_{12}\} \in U_{sqp}$. Once again, we take e_1 to be in $\mathcal{S}$. Then there is some basic open set $U_{11} \times U_{12} / \sim$, centered at e_1 , that lies inside U_{sqp} . Let, B_{11} (centered at v_{11}) and B_{12} (centered at v_{12}) be open balls in $\mathcal{V}$ of equal radius, r , contained within U_{11} and U_{12} respectively.

Let $e_2 = \{v_{21}, v_{22}\} \in \mathcal{E}$ with $d_{\mathcal{E}}(e_1, e_2) < r$ (which of course means that e_2 is also in $\mathcal{S}$). Then

$$\begin{aligned} r &> d_{\mathcal{E}}(e_1, e_2) \\ &= d_{(\mathcal{V} \times \mathcal{V}) / \sim}(e_1, e_2) \\ &= \min_{i \neq j \in \{1,2\}} d_{\mathcal{V} \times \mathcal{V}}((v_{11}, v_{12}), (v_{2i}, v_{2j})) \\ &= \min_{i \neq j \in \{1,2\}} \max\{d_{\mathcal{V}}(v_{11}, v_{2i}), d_{\mathcal{V}}(v_{12}, v_{2j})\}. \end{aligned}$$

That means that either $v_{21} \in B_{11}$ and $v_{22} \in B_{12}$ or the other way around (or both). Either way, $e_2 \in B_{11} \times B_{12} / \sim$, so there is an open $d_{\mathcal{E}}$ ball completely contained within U_{sqp} around every point in U_{sqp} . Thus U_{sqp} is open in $\mathcal{T}_{\mathcal{E}}$. $\square$

That's good news because it wasn't at all obvious that we could simply take two open sets on $\mathcal{V}$ and know that the edges whose ends fall into those open sets themselves form open sets on $\mathcal{S}$ and $\mathcal{E}$.

3.2.3 COMPACTNESS OF $\mathcal{V}$

Now that we have topologies on $\mathcal{V}$ and $\mathcal{E}$, let's talk for a moment about $\mathcal{V}$. Consider the tensegrity shown in [Figure 3.2.3](#). This tensegrity is our (bar equivalent) crossed-

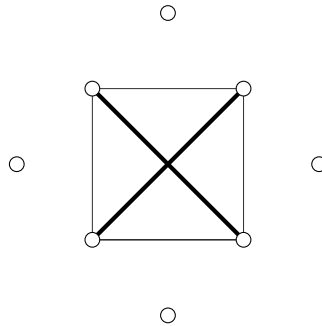


Figure 3.2.3: A tensegrity with “free vertices”.

square plus four extra vertices. What effect do those extra vertices have on the bar equivalence of this tensegrity? None whatsoever.

Claim 3.2.9. *Adding or removing vertices which have no edges connected to them does not change the $\mathcal{X}$ -bar equivalence of a tensegrity.*

Proof. $\mathcal{X}$ -bar equivalence and partial $\mathcal{X}$ -bar equivalence are conditions about $Y(\mathcal{X})$. Adding or removing vertices, so long as it does not change the edgeset, cannot change $Y(\mathcal{X})$. □

So, when it is convenient, we will consider $\mathcal{V}$ to have only vertices to which edges are attached. That will give us an additional attribute for $\mathcal{V}$.

Claim 3.2.10. *If $\mathcal{V}$ consists only of vertices to which edges in $\mathcal{E}$ are attached, then $\mathcal{V}$ is compact.*

Proof. Suppose every vertex in $\mathcal{V}$ has an edge attached to it. Let $\{U_\alpha\}$ be an open cover of $\mathcal{V}$. Then, by the work we've done above, the sets $U_{\alpha_1} \times U_{\alpha_2} / \sim \cap \mathcal{E}$ (where U_{α_1} and U_{α_2} are in $\{U_\alpha\}$) are open sets on $\mathcal{S}$ and $\mathcal{C}$. Since every edge has to connect to

vertices at both ends, these open sets cover $\mathcal{E}$. Since $\mathcal{E}$ is compact, a finite number of them, $U_{\alpha_{11}} \times U_{\alpha_{21}} / \sim, \dots, U_{\alpha_{1n}} \times U_{\alpha_{2n}} / \sim$ will suffice to cover $\mathcal{E}$.

We want to show that the open sets $U_{\alpha_{11}}, \dots, U_{\alpha_{1n}}, U_{\alpha_{21}}, \dots, U_{\alpha_{2n}}$ cover $\mathcal{V}$. Suppose $v_1 \in \mathcal{V}$. Then, since all vertices of $\mathcal{V}$ have edges connected, there must be some $e = \{v_1, v_2\}$ in $\mathcal{E}$. But then there must be some $U_{\alpha_{1i}} \times U_{\alpha_{2i}} / \sim$ which contains e . So either $v_1 \in U_{\alpha_{1i}}$ or $v_1 \in U_{\alpha_{2i}}$.

So the $U_{\alpha_{11}}, \dots, U_{\alpha_{1n}}, U_{\alpha_{21}}, \dots, U_{\alpha_{2n}}$ cover $\mathcal{V}$. Hence $\mathcal{V}$ is compact. $\square$

Thus, while we never require that $\mathcal{V}$ be compact, we discover that if $\mathcal{E}$ is compact, then $\mathcal{V}$ is effectively compact as well.

3.2.4 THE NEW TOPOLOGIES MADE VISIBLE

Let's take a moment to look at a specific example. Suppose we have the tensegrity shown in [Figure 3.2.4](#) on the following page(b), whose the vertex curve is shown in [Figure 3.2.4\(a\)](#). The vertex curve consists of two connected components: a central segment running from $(0, 0)$ to $(0, 2)$ and a stadium curve running at distance 2 around that segment. Struts are connected orthogonally to the curve. Cables are connected in the piecewise linear fashion shown in [Figure 3.2.5](#) on page 50 (essentially, endpoints on the curves move with one fixed speed and endpoints on the straight segments with another).

In [Figure 3.2.5](#), we show the sets $\mathcal{S}$ and $\mathcal{C}$ for the stadium curve tensegrity. These two sets are subsets of the space $(\mathcal{V} \times \mathcal{V}) / \sim$. If we plot $\mathcal{V} \times \mathcal{V}$ as a square, the space $(\mathcal{V} \times \mathcal{V}) / \sim$ can be viewed as either the triangle above the diagonal or the one below. We've used one of each for the sake of space.

The parameter for $\mathcal{V}$, as it runs along the edge of the square, first sweeps out the component whose vertices lie along the center segment (from bottom to top) and then covers the outer curve, starting at the point marked in [Figure 3.2.4](#) and running in the direction of the arrow around the outside.

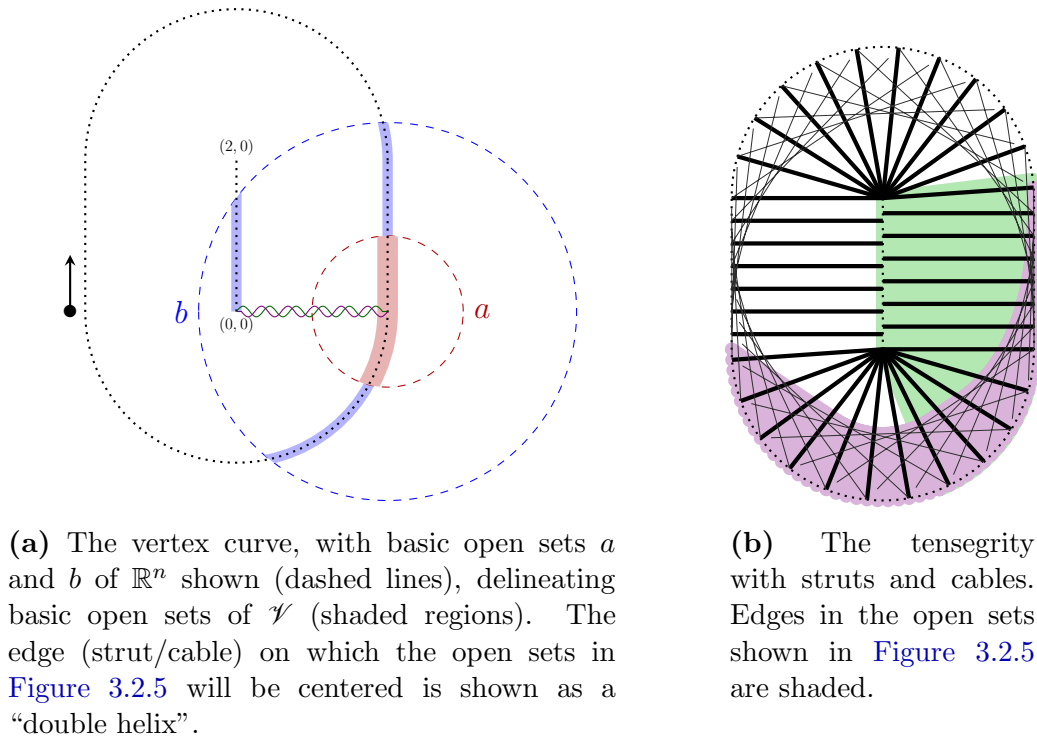


Figure 3.2.4: The stadium-curve tensegrity.

The open sets a and b of $\mathcal{V}$ are shown marked along the edges of the square. The struts are plotted in the upper triangle, while the cables are shown in the lower one. We’ve marked two basic open sets (in both cases, a ball of radius $\sqrt{5}$ around the edge shown in Figure 3.2.4(a)).

3.2.5 ABOUT NON-COMPACT EDGESETS

Let’s talk again for a moment about the example from Figure 3.2.1 on page 42. Why doesn’t it trouble us to exclude this case? Frankly, because it is indistinguishable from the case shown in Figure 3.2.6 on page 51. Here, the strut at 0 and the cable at π have

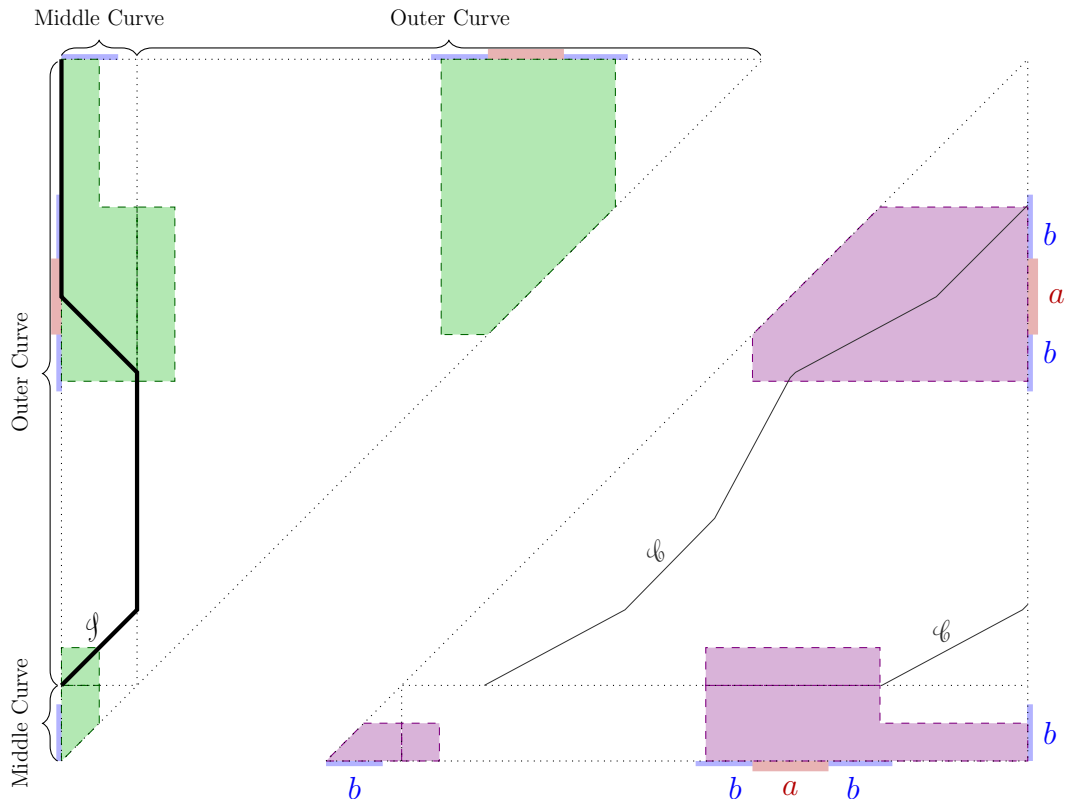


Figure 3.2.5: Here we show the sets $\mathcal{S}$ and $\mathcal{C}$ for the stadium-curve tensegrity of Figure 3.2.4 on the previous page with one basic open set for each connected component of $\mathcal{C}$ (shaded, with dashed outlines). See the text for a more complete description of this figure.

been replaced with bars. Now, both $\mathcal{S}$ and $\mathcal{C}$ are compact⁴ and all of our results hold.

This tensegrity, however, works exactly as the one in Figure 3.2.1 did, because the elements of $\text{VF}(\mathcal{V})$ are continuous. Any vector field which expanded the strut at 0 , s_0 , would have had to expand all edges in an open set around s_0 . But any such open set contains cables, which cannot expand.

Thus s_0 cannot change length and it might as well be a bar. By the same type of argument, the cable at π is effectively a bar as well.

⁴As we noted in the footnote on page 41, $\mathcal{V} \times \mathcal{V}$ is compact and since $\mathcal{S}$ and $\mathcal{C}$ are now closed, their preimages in $\mathcal{V} \times \mathcal{V}$ are compact and hence they are themselves compact.

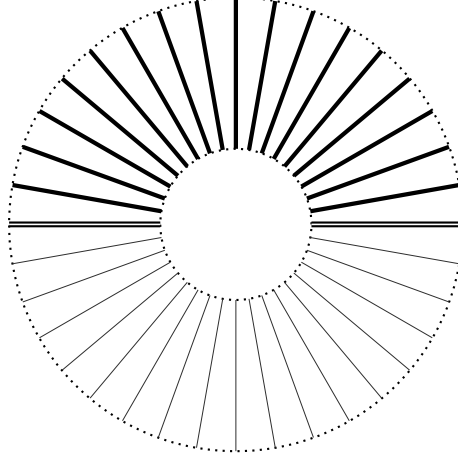


Figure 3.2.6: The tensegrity of Figure 3.2.1 “compactified”. Bars have been added at 0 and π . Now $\mathcal{S}$ and $\mathcal{C}$ are both compact.

Proposition 3.2.11. *For any $V \in \text{VF}(\mathcal{V})$, $YV \in C(\mathcal{E})^+ \setminus \{\mathbf{0}\}$ if and only if $Y_c V \in C(\mathcal{E}_c)^+ \setminus \{\mathbf{0}\}$, where $\mathcal{E}_c$ is the disjoint union of the closures of $\mathcal{S}$ and $\mathcal{C}$ (that is, $\mathcal{E}_c = \overline{\mathcal{S}} \sqcup \overline{\mathcal{C}}$), and Y_c is the rigidity operator which results from replacing $\mathcal{E}$ with $\mathcal{E}_c$.*

Proof. If $Y_c V \in C(\mathcal{E}_c)^+$, then YV , which is the restriction of $Y_c V$ to $\mathcal{E}$, must be in $C(\mathcal{E})^+$. If $Y_c V \neq \mathbf{0}$, then there must be some edge $e_c \in \mathcal{E}_c$ for which $Y_c V(e_c) > 0$. But because $Y_c V$ is continuous, there is some open set on which $Y_c V > 0$, and that open set must intersect $\mathcal{E}$ (as e_c is in the closure of either $\mathcal{S}$ or $\mathcal{C}$). So YV must also be nonzero. That’s one direction.

Conversely, assume that $YV \in C(\mathcal{E})^+$. Then its image must be contained in the interval $[0, \infty)$. We want to show that the $\text{im } Y_c V$ is also contained in that interval. Suppose, to the contrary, there is some edge e_n such that $Y_c V(e_n) < 0$. Now e_n is either in $\overline{\mathcal{S}}$ or in $\overline{\mathcal{C}}$. Without loss of generality, take it to be in $\overline{\mathcal{S}}$.

Since $Y_c V$ is continuous, there must be an open set U_- containing e_n on which $Y_c V$ is strictly negative. But since $e_n \in \overline{\mathcal{S}}$, we must have $U_- \cap \mathcal{S} \neq \emptyset$, and $Y_c \geq 0$ on $\mathcal{S}$, so it cannot be strictly negative on U_- . That contradiction shows that $\text{im } Y_c V \in [0, \infty)$.

Also, if $YV \neq \mathbf{0}$, then on some $e \in \mathcal{E}$, $YV(e) \neq 0$ and hence $Y_c V(e) \neq \mathbf{0}$, so $Y_c V \neq \mathbf{0}$. $\square$

3.2.6 DISCRETE TOPOLOGIES

Just to finish things off well, let's end this section by showing that in the finite case, the topologies on $\mathcal{V}$ and $\mathcal{E}$ are discrete.

Proposition 3.2.12. *If $\mathcal{V}$ and $\mathcal{E}$ are finite sets, the topologies on them are discrete.*

Proof. If $\mathcal{V}$ is a finite set, then $p(\mathcal{V})$ is a finite set of points in $\mathbb{R}^n$. Since p is 1-1, if $v_1 \neq v_2$, then $d_{\mathcal{V}}(v_1, v_2) = \|p(v_1) - p(v_2)\| > 0$, so open balls in $\mathcal{V}$ centered on v_1 and v_2 and having radius less than half of $d_{\mathcal{V}}(v_1, v_2)$ separate v_1 and v_2 . For any given vertex, v , the (finite) intersection of all such balls is an open set which contains v and no other element of $\mathcal{V}$. Thus we have the discrete topology on $\mathcal{V}$.

Similarly, if $e_1, e_2 \in \mathcal{E}$, then either they are both struts (or both cables) or there is one of each. Since the sets $\mathbf{S}$, $\mathbf{C}$ and $\mathbf{B}$ are pairwise disjoint, no two struts connect the same pair of vertices and no two cables connect the same pair of vertices. So if e_1 and e_2 are of the same type, they are some positive distance apart. But if they are of different type, that is even more true. Using the same process as we did for $\mathcal{V}$, we see that this topology is also discrete. $\square$

3.3 $C(\mathcal{E})$ AND $C^*(\mathcal{E})$

Since $C(\mathcal{E})$ and $C^*(\mathcal{E})$ are no longer finite vector spaces, the Euclidean norm will no longer serve. We'll put the sup norm on $C(\mathcal{E})$ and the operator norm on $C^*(\mathcal{E})$. That is, for $f \in C(\mathcal{E})$ and $\mu \in C^*(\mathcal{E})$, we have

$$\|f\| = \sup_{e \in \mathcal{E}} |f(e)| \text{ and}$$

$$\|\mu\| = \sup_{\|f\|=1} |\mu f|.$$

We can induce a metric on each by considering $\|x - y\|$ to be the distance between x and y and establish on each the associated metric topology.

One quick definition:

Definition 3.3.1 (see, for example, [Folland \(1999, p. 132\)](#)). If X is a topological space, the *support* of a function $f: X \rightarrow \mathbb{R}$, denoted $\text{supp } f$, is the smallest closed set $S \subset X$ such that f is zero everywhere on the complement $X \setminus S$. In other words, S is the closure of the set of all points $x \in X$ for which $f(x) \neq 0$.

And now we're prepared to talk about the interiors of $C(\mathcal{E})^+$ and $C(\mathcal{E})^-$.

Lemma 3.3.2. *The interiors of the nonnegative and nonpositive orthants of $C(\mathcal{E})$, that is $\text{int } C(\mathcal{E})^+$ and $\text{int } C(\mathcal{E})^-$, consist of the strictly positive and strictly negative continuous functions from $\mathcal{E}$ to $\mathbb{R}$, respectively.*

Proof. We will prove here that $\text{int } C(\mathcal{E})^+$ consists of the strictly positive elements of $C(\mathcal{E})$. The proof for $\text{int } C(\mathcal{E})^-$ is identical except for sign. We need to establish our result using the Euclidean norm for the finite case and the sup norm for the general case.

Let f be a strictly positive, continuous function on $\mathcal{E}$. Then, since $\mathcal{E}$ is compact, f achieves a (strictly positive) minimum on $\mathcal{E}$. Now let $g \in C(\mathcal{E})$ with $\|f - g\| = d < \min f$. Of course g also achieves a minimum on $\mathcal{E}$ at some edge e . Then, under either norm,

$$\min g = g(e) = f(e) - (f(e) - g(e)) \geq f(e) - \|f - g\| \geq \min f - d > 0,$$

so g is strictly positive. Thus f is in $\text{int } C(\mathcal{E})^+$.

Conversely, let $f \in \text{int } C(\mathcal{E})^+$. Then there exists some $\delta > 0$ such that $g \in C(\mathcal{E})^+$ for all $\|f - g\| < \delta$. Let $e \in \mathcal{E}$ be where f achieves its minimum. Suppose $f(e) < \delta$.

In the finite case, let $g = f$ everywhere except at e and set $g(e) = \frac{f(e) - \delta}{2}$. In the infinite case, select an open set U containing e and let $g_e \in C(\mathcal{E})^+$ with $\sup g_e =$

$g_e(e) = \frac{f(e)-\delta}{2}$ and $\text{supp } g_e \subset U$, and set $g = f - g_e$. Then in both cases, we have $\|f - g\| = \frac{f(e)+\delta}{2} \in (0, \delta)$, but $g(e) = \frac{f(e)-\delta}{2} < 0$. That contradiction means that we must have $f(e) = \min f \geq \delta > 0$, and f is strictly positive. $\square$

Next we turn our attention from $C(\mathcal{E})$ to its topological dual, $C^*(\mathcal{E})$, which, we remember, consists of the continuous linear functionals on $C(\mathcal{E})$ (Definition 2.1.13 on page 22). We'll want to relate the elements of $C^*(\mathcal{E})$ not only to $C(\mathcal{E})$ but also to $\mathcal{E}$ itself.

First, we should note that for linear functionals, “continuous” (Definition 3.7.4 on page 76) and “bounded” (Definition 3.4.1 on page 60) are connected:

Theorem 3.3.3. *For a linear transformation Λ of a normed linear space X into a normed linear space Z , each of the following three conditions implies the other two:*

- (a) Λ is bounded.
- (b) Λ is continuous.
- (c) Λ is continuous at one point of X .

Proof. See Rudin (1987, p. 96). $\square$

Next, a quick definition.

Definition 3.3.4 (Rudin (1987, p. 47)). A measure is called *regular* if, for any measurable set E ,

$$\mu(E) = \inf\{\mu(U) : E \subset U, U \text{ open}\}$$

and

$$\mu(E) = \sup\{\mu(K) : K \subset E, K \text{ compact}\}.$$

And now we're ready for a theorem that will allow us to understand the elements of $C^*(\mathcal{E})$ not only as continuous linear functionals on $C(\mathcal{E})$, but also as measures on $\mathcal{E}$.

Theorem 3.3.5 (Riesz Representation Theorem). *If X is a compact Hausdorff space, then every bounded linear functional Φ on $C(X)$ is represented by a unique regular Borel measure μ , in the sense that*

$$\Phi f = \int_X f d\mu$$

for every $f \in C(X)$.

Proof. See [Rudin \(1987, p. 130\)](#). Rudin actually gives the theorem for locally compact spaces and requires the functions to be in $C_0(X)$ (the continuous functions which vanish at infinity), but for a compact set, $C_0(X) = C(X)$. Also, he returns a regular complex measure ν , but our $\mu = |\nu|$, the total variation of ν , which is regular since ν is regular ([Rudin, 1987](#), pp. 70, 116–117, 130). $\square$

Since we are dealing in measures, we'll want to talk about “strictly positive measures” and “semipositive measures”, and since the Riesz Representation theorem gives it to us, we'll expect regularity. We'll take the following definitions.

Definition 3.3.6. A *strictly positive measure* is a regular measure μ on a set X such that for every nonempty, open $U \subset X$, $\mu(U) > 0$. A *semipositive measure* is a nonzero regular measure μ on a set X such that for every nonempty, open $U \subset X$, $\mu(U) \geq 0$.

Following general practice, we will use the term *positive measure* to refer to what might more accurately be called “nonnegative measure”. The regular positive measures differ from the semipositive measures only in that the zero measure is considered a regular positive measure.

It seems now that we have competing definitions for “strictly positive” and “semipositive” for those measures that are stresses of our tensegrity (in [Definition 2.1.16](#) on [page 22](#) they were defined in terms of $\mu(f)$ rather than $\mu(U)$). In [Proposition 3.3.17](#) on [page 58](#), we'll show that the definitions are equivalent, but we have some work to do before we get there.

First, we need to define the support of a measure.

Definition 3.3.7 (see, for example, Morrison (2001, p. 152)). If μ is a measure on a set X (with its associated σ -algebra), then the *support* of μ , denoted $\text{supp } \mu$ is the set of $x \in X$ such that $|\mu(U)| > 0$ for all open neighborhoods U of x .

There are a few lemmas about supports we will want. The proof for the first was found on Wikipedia (2006).

Lemma 3.3.8. *The support of a measure μ on a topological set X is a closed set.*

Proof. Let μ be a measure on a topological set X and let $\{x_i \in \text{supp } \mu\}$ be a sequence which approaches some limit $x \in X$. If $U \subset X$ is any open set which contains x , then, since x is the limit of the x_i , U must intersect $\{x_i\}$. But that means U is an open set containing some element of $\text{supp } \mu$ and hence $|\mu(U)| > 0$. So $x \in \text{supp } \mu$. $\square$

Lemma 3.3.9. *If μ is a regular positive measure on a topological set X and A is an open subset of X such that $\text{supp } \mu \cap A = \emptyset$, then $\mu(A) = 0$.*

Proof. We note first that if $x \in A$, then $x \notin \text{supp } \mu$, so there must be some open neighborhood U_x of x with measure 0. If K is any compact subset of A , then K can be covered by a finite number of these U_x . So $0 \leq \mu(K) \leq \sum_{x \in K} \mu(U_x) = 0$. But since μ is regular, $\mu(A) = \sup\{\mu(K) : K \subset A, K \text{ compact}\}$, so $\mu(A) = 0$. $\square$

Now we can show how positivity relates to the support of a measure.

Lemma 3.3.10. *μ is a strictly positive measure on $\mathcal{E}$ if and only if μ is a regular positive measure with $\text{supp } \mu = \mathcal{E}$.*

Proof. Let μ be strictly positive and let $e \in \mathcal{E}$. Then, for every neighborhood U of e , $\mu(U) > 0$ (after all, μ is strictly positive), so $e \in \text{supp } \mu$.

Conversely, let μ be a regular positive measure on $\mathcal{E}$ with $\text{supp } \mu = \mathcal{E}$ and let $U \subset \mathcal{E}$ be nonempty and open. Then $U \cap \text{supp } \mu = U \neq \emptyset$, so $\mu(U) > 0$. $\square$

Lemma 3.3.11. μ is a semipositive measure on $\mathcal{E}$ if and only μ is a regular positive measure with $\text{supp } \mu \neq \emptyset$.

Proof. Let μ be semipositive. Suppose, to the contrary, that $\text{supp } \mu = \emptyset$. Then, $\mathcal{E}$ is an open set disjoint from $\text{supp } \mu$, so by Lemma 3.3.9, $\mu(\mathcal{E}) = 0$ and since μ is semipositive, we must have $\mu(U) = 0$ for every $U \subset \mathcal{E}$. This contradiction shows that $\text{supp } \mu \neq \emptyset$.

Conversely, let μ be a regular positive measure on $\mathcal{E}$ with $\text{supp } \mu \neq \emptyset$ and let $e \in \text{supp } \mu$. Then, since μ is positive, $\mu(U) \geq 0$ for all $U \subset \mathcal{E}$. Furthermore, $\mathcal{E}$ is an open set which contains e and hence $\mu(\mathcal{E}) > 0$. $\square$

We'll need to know about the Jordan decomposition of a measure, which allows a signed measure to be “split” into positive and negative parts and the Urysohn Lemma, which allows us to create certain continuous maps. Here are the details:

Definition 3.3.12 (Folland (1999, p. 87)). Two signed measures μ and ν on a measurable space $(X, \mathcal{M})$ are *mutually singular* if there exist $E, F \in \mathcal{M}$ such that $E \cap F = \emptyset$, $E \cup F = X$, E is null for μ and F is null for ν . We denote this $\mu \perp \nu$.

Theorem 3.3.13 (Jordan Decomposition Theorem). *If ν is a signed measure, there exist unique positive measures ν^+ and ν^- such that $\nu = \nu^+ - \nu^-$ and $\nu^+ \perp \nu^-$.*

Proof. See Folland (1999, p. 87). $\square$

Definition 3.3.14 (Munkres (2000, p.195)). A space X is said to be *normal* if, for each pair A, B of disjoint closed sets of X , there exist disjoint open sets containing A and B , respectively.

Lemma 3.3.15. *Every metric space is normal.*

Proof. For this proof, we'll take some guidance from a hint in Lawson (2003, p. 52). We let X be a metric space with metric d and let A and B be disjoint closed subsets

of X . For each $a \in A$, we know that since B is closed and A and B are disjoint, a is not in the closure of B . So there must be some strictly positive infimal distance d_a between a and all of B . Let U_a be the open ball of radius $d_a/2$ centered at a . Clearly, $U_a \cap B = \emptyset$. Furthermore, the collection of all such U_a covers A . Similarly, we can cover B with U_b where the radius of a given U_b is half the minimum distance between b and all of A .

Now let $U_A = \bigcup U_a$ and $U_B = \bigcup U_b$. We'd like to show that $U_A \cap U_B = \emptyset$. Suppose, to the contrary, that some element $x \in X$ lies in both U_A and U_B . Since x is in a union of sets, it must be in one of the sets. That is, there must be some $a \in A$ and some $b \in B$ such that $x \in U_a \cap U_b$. Let d_a be the infimal distance between a and B and d_b the infimal distance between b and A . Then the distance between the two points $d(a, b)$ has to be at least as large as d_a and d_b . But by the triangle inequality

$$\max\{d_a, d_b\} \leq d(a, b) \leq d(a, x) + d(x, b) < \frac{d_a}{2} + \frac{d_b}{2} \leq 2 \frac{\max\{d_a, d_b\}}{2} = \max\{d_a, d_b\}.$$

That contradiction shows us that there can be no such x and hence U_A and U_B are disjoint. So every metric space is normal. $\square$

Theorem 3.3.16 (Urysohn Lemma). *Let X be a normal space; let A and B be disjoint closed subsets of X . Let $[a, b]$ be a closed interval in the real line. Then there exists a continuous map $f: X \rightarrow [a, b]$ such that $f(x) = a$ for every x in A , and $f(x) = b$ for every x in B .*

Proof. See, for example, [Munkres \(2000, p. 207\)](#). $\square$

Proposition 3.3.17. *Let $\mu \in C^*(\mathcal{E})$ with $\mu \neq \mathbf{0}$. Then $\mu(f) > 0$ (resp. $\mu(f) \geq 0$) for all nonzero $f \in C(\mathcal{E})^+$ if and only if $\mu(U) > 0$ (resp. $\mu(U) \geq 0$) for all nonempty, open $U \subset \mathcal{E}$.*

Proof. We'll start by showing that if $\mu(f) \geq 0$ for all $f \in C(\mathcal{E})^+$, then $\mu(U) \geq 0$ for all open $U \subset E$.

Let $\mu \in C^*(\mathcal{E})^+$ such that $\mu \neq \mathbf{0}$ and $\mu(f) \geq 0$ for all $f \in C(\mathcal{E})^+$. Let $\mu = \mu^+ - \mu^-$ be the Jordan decomposition of μ . By Lemma 3.3.8, $\text{supp } \mu^+$ is closed, so $U_s = \mathcal{E} \setminus \text{supp } \mu^+$ is open. If $U_s = \emptyset$, then by Lemma 3.3.10, μ^+ is a strictly positive measure. Since $\mu^+ \perp \mu^-$ in the Jordan decompositions, then the set F from Definition 3.3.12 must be the empty set. So μ^- is the zero measure and μ is strictly positive as well and we are done.

Suppose, then, that $U_s \neq \emptyset$ and, further, that there exists some $U_- \subset \mathcal{E}$ such that $\mu(U_-) < 0$. If U_- were completely contained in $\text{supp } \mu^+$, we'd have $\mu^+(U_-) > 0$ and $\mu^-(U_-) = 0$ and thus $\mu(U_-) > 0$. Since that isn't the case, we know that U_- reaches (or lies completely) outside of $\text{supp } \mu^+$. That is, $U_i = U_- \cap U_s \neq \emptyset$ and $\mu(U_i) = -\mu^-(U_i) < 0$.

Since $-\mu^-(U_i) < 0$, U_i must intersect the support of μ^- nontrivially (by Lemma 3.3.9 on page 56), so define $E_i = \text{supp } \mu^- \cap \overline{U_i}$, which means that $\mu(E_i) < 0$. Now E_i and $\text{supp } \mu^+$ are disjoint closed sets in the metric (hence normal, by Lemma 3.3.15) space $\mathcal{E}$, so by the Urysohn Lemma, there is a continuous function f which takes the value 1 on E_i and the value 0 on $\text{supp } \mu^+$ and whose values lie in the interval $[0, 1]$ elsewhere. Clearly $f \in C(\mathcal{E})^+$.

So now we have $\mu(f) = \int_{\mathcal{E}} f d\mu = \int_{\text{supp } f} f d\mu$. But by design, $\text{supp } f \cap \text{supp } \mu^+ = \emptyset$, so we have $\mu(f) = -\int_{\text{supp } f} f d\mu^- \leq -\int_{E_i} f d\mu^- = \mu^-(E_i) < 0$.

That contradicts our hypothesis that $\mu(f) \geq 0$ for all $f \in C(\mathcal{E})^+$, so there can be no such U_- . Hence $\mu(U) \geq 0$ for all $U \subset \mathcal{E}$.

Next we show that if $\mu(f) > 0$ for all nonzero $f \in C(\mathcal{E})^+$, then $\mu(U) > 0$ for all nonempty, open $U \subset \mathcal{E}$.

Let $\mu \in C^*(\mathcal{E})^+$ such that $\mu(f) > 0$ for all nonzero $f \in C(\mathcal{E})^+$. Now we know (by the work we just finished) that $\mu(U) \geq 0$ for all nonempty, open $U \subset \mathcal{E}$, but suppose that for one such set, $\mu(U) = 0$. Select any nonnegative, continuous function f_U on $\mathcal{E}$ with $|f_U| = 1$ and $\text{supp } f_U \subset U$.

Then, defining $\mathbf{1}$ to be the constant function $\mathbf{1}(e) \equiv 1$ on $\mathcal{E}$, we have

$$0 < \mu(f_U) = \int_{\mathcal{E}} f_U d\mu = \int_{\text{supp } f_U} f_U d\mu \leq \int_{\text{supp } f_U} \mathbf{1} d\mu \leq \mu(U) = 0.$$

That contradiction show that there can be no such U .

Finally, we work the other direction. The material given here is for the strict inequalities, but it works *a fortiori* if they are not strict.

Let $\mu \in C^*(\mathcal{E})^+$ such that $\mu(U) > 0$ for all nonempty, open $U \subset \mathcal{E}$ and let f be a nonzero element of $C(\mathcal{E})^+$. Since f is nonzero, there must be some $e_+ \in \mathcal{E}$ such that $f(e_+) > 0$. As f is continuous $U_+ = f^{-1}((f(e_+)/2, \infty))$ is nonempty and open. Thus

$$\mu(f) = \int_{\text{supp } f} f d\mu \geq \int_{U_+} f d\mu \geq \int_{U_+} \frac{f(e_+)}{2} d\mu = \frac{f(e_+)}{2} \mu(U_+) > 0.$$

□

3.4 VF($\mathcal{V}$) AND Y

$\text{VF}(\mathcal{V})$ consists of the continuous vector fields from $\mathcal{V}$ to $\mathbb{R}^n$. We'll want it to be a topological space, so we'll give it the sup norm. That is, for $V \in \text{VF}(\mathcal{V})$,

$$\|V\| = \sup_{v \in \mathcal{V}} \|V(v)\|.$$

Then we'll use the norm to give us a metric and endow $\text{VF}(\mathcal{V})$ with the associated metric topology.

When we first met Y , back on page 18, we defined it as a map from $\text{VF}(\mathcal{V})$ to $C(\mathcal{E})$, but much of what we have done since then has thought of Y simply as a matrix. That was fine for the finite world, but we need to get back to thinking of Y as a map. As a map, Y has some nice attributes.

Definition 3.4.1 (Rudin (1987, p. 96)). Consider a linear transformation Λ from a normed linear space X into a normed linear space Z , and define its *norm* by

$$\|\Lambda\| = \sup\{\|\Lambda x\| : x \in X, \|x\| \leq 1\}.$$

If $\|\Lambda\| < \infty$, then Λ is called a *bounded linear transformation*.

In a moment we'll see that our map $\|Y\|$ is bounded. First, though, we have a lemma about edge lengths.

Lemma 3.4.2. *The edge length $\|p(v_1) - p(v_2)\|$ is continuous in $e \in \mathcal{E}$.*

Proof. Let $\varepsilon > 0$. Select $\delta = \varepsilon/2$. Then, whenever $e_1 = \{v_{11}, v_{12}\}, e_2 = \{v_{21}, v_{22}\} \in \mathcal{E}$ with $d_{\mathcal{E}}(e_1, e_2) < \delta$, we have either

$$\|p(v_{11}) - p(v_{21})\| < \delta \text{ and } \|p(v_{12}) - p(v_{22})\| < \delta \quad (3.2)$$

or

$$\|p(v_{11}) - p(v_{22})\| < \delta \text{ and } \|p(v_{12}) - p(v_{21})\| < \delta. \quad (3.3)$$

But then

$$\begin{aligned} \left| \|p(v_{11}) - p(v_{12})\| - \|p(v_{21}) - p(v_{22})\| \right| &\leq \left| \|p(v_{11}) - p(v_{12}) - p(v_{21}) + p(v_{22})\| \right| \\ &\leq \|p(v_{11}) - p(v_{21})\| + \|p(v_{22}) - p(v_{12})\| \end{aligned} \quad (3.4)$$

and also

$$\begin{aligned} \left| \|p(v_{11}) - p(v_{12})\| - \|p(v_{21}) - p(v_{22})\| \right| &= \left| \|p(v_{11}) - p(v_{12})\| - \|p(v_{22}) - p(v_{21})\| \right| \\ &\leq \left| \|p(v_{11}) - p(v_{12}) - p(v_{22}) + p(v_{21})\| \right| \\ &\leq \|p(v_{11}) - p(v_{22})\| + \|p(v_{21}) - p(v_{12})\|. \end{aligned} \quad (3.5)$$

Since both Equation (3.4) and Equation (3.5) are true, and either Equation (3.2) or Equation (3.3) is true as well, we have

$$\left| \|p(v_{11}) - p(v_{12})\| - \|p(v_{21}) - p(v_{22})\| \right| \leq \delta + \delta = \varepsilon$$

and we are done. □

Proposition 3.4.3. *For any rigidity operator Y , $\|Y\| < \infty$.*

Proof. Let $V \in \text{VF}(\mathcal{V})$ with $\|V\| \leq 1$. Then $\sup_{v \in \mathcal{V}} \|V(v)\| \leq 1$. So for any given edge, $e = \{v_1, v_2\}$, we have $\|V(v_1) - V(v_2)\| \leq 2$. By Lemma 3.4.2, edge length is a continuous function on $\mathcal{E}$, and $\mathcal{E}$ is compact, so there is some edge e_L of longest length $L < \infty$. Thus, whenever $\|V\| \leq 1$, we have

$$\|YV\| = \sup_{\{v_1, v_2\} \in \mathcal{E}} |(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2))| \leq 2L.$$

Since $\|Y\|$ is defined as

$$\|Y\| = \sup\{\|YV\| : V \in \text{VF}(\mathcal{V}), \|V\| \leq 1\},$$

we have $\|Y\| \leq 2L < \infty$. □

We'd like to show that Y is a continuous map from $\text{VF}(\mathcal{V})$ to $C(\mathcal{E})$. To do that we'll need a little machinery. We need a theorem about continuity in quotient spaces and a lemma about a continuous function on $\mathcal{V} \times \mathcal{V}$.

Theorem 3.4.4. *Let $\pi: X \rightarrow W$ be a quotient map. Let Z be a space and let $g: X \rightarrow Z$ be a map that is constant on each set $\pi^{-1}(w)$, for $w \in W$. Then g induces a map $f: W \rightarrow Z$ such that $f \circ \pi = g$. The induced map f is continuous if and only if g is continuous.*

$$\begin{array}{ccc} X & & \\ \pi \downarrow & \searrow g & \\ W & \xrightarrow{\quad f \quad} & Z \end{array}$$

Proof. See Munkres (2000, p. 142). □

Lemma 3.4.5. *If $f(v)$ is a continuous function from $\mathcal{V}$ to $\mathbb{R}^n$, then $f(v_1) - f(v_2)$ is a continuous function from $\mathcal{V} \times \mathcal{V}$ to $\mathbb{R}^n$.*

Proof. Let $\varepsilon > 0$. Since $f(v)$ is continuous, there exists some $\delta > 0$ such that whenever v_1 and v_2 are in $\mathcal{V}$ with $d_{\mathcal{V}}(v_1, v_2) < \delta$, we have $\|f(v_1) - f(v_2)\| < \varepsilon/2$. Now,

whenever (v_1, v_2) and (v_3, v_4) are in $\mathcal{V} \times \mathcal{V}$ with $d_{\mathcal{V} \times \mathcal{V}}((v_1, v_2), (v_3, v_4)) < \delta$, we have $d_{\mathcal{V}}(v_1, v_3) < \delta$ and $d_{\mathcal{V}}(v_2, v_4) < \delta$. But then

$$\begin{aligned} \|(f(v_1) - f(v_2)) - (f(v_3) - f(v_4))\| &= \|f(v_1) - f(v_2) - f(v_3) + f(v_4)\| \\ &= \|(f(v_1) - f(v_3)) + (f(v_4) - f(v_2))\| \\ &\leq \|f(v_1) - f(v_3)\| + \|f(v_4) - f(v_2)\| \\ &< \frac{\varepsilon}{2} + \frac{\varepsilon}{2} = \varepsilon. \end{aligned}$$

□

Now we can state and prove our proposition.

Proposition 3.4.6. *The rigidity operator Y is a continuous linear map from $\mathbf{VF}(\mathcal{V})$ to $C(\mathcal{E})$.*

Proof. We have two things to accomplish. First, we need to show that continuous vector fields get mapped to continuous functions under Y . Second, we need to show that Y is a continuous linear map.

Remembering that YV is the pointwise dot product of the functions $V(v_1) - V(v_2)$ and $p(v_1) - p(v_2)$, we note that

$$(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2)) = (V(v_2) - V(v_1)) \cdot (p(v_2) - p(v_1)).$$

So if π is the quotient map from $\mathcal{V} \times \mathcal{V} \rightarrow (\mathcal{V} \times \mathcal{V})/_\sim$, we see that $(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2))$ is constant on $\pi^{-1}(\{v_1, v_2\})$.

Now V is continuous, since it is an element of $\mathbf{VF}(\mathcal{V})$, and p is continuous because the topology on $\mathcal{V}$ is defined by it. So by Lemma 3.4.5, $V(v_1) - V(v_2)$ and $p(v_1) - p(v_2)$ are both continuous on $\mathcal{V} \times \mathcal{V}$. Hence $(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2))$ is continuous on $\mathcal{V} \times \mathcal{V}$. But by Theorem 3.4.4, then, $YV = (V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2))$ is continuous on $(\mathcal{V} \times \mathcal{V})/_\sim$ and thus on $\mathcal{E}$.

So we've shown that Y maps from $\text{VF}(\mathcal{V})$ to $C(\mathcal{E})$. Let's show that it does so in a continuous fashion. We could do that by invoking Theorem 3.3.3 on page 54, but doing the work ourselves is not difficult.

Let U be an open set of functions in $C(\mathcal{E})$. If $\text{im } Y \cap U = \emptyset$, then $Y^{-1}U = \emptyset$, which is open. Otherwise, let $V_1 \in Y^{-1}U$. Now there is a basic open set B around YV_1 contained in U with radius $r > 0$. Let $V_2 \in \text{VF}(\mathcal{V})$ with $\|V_1 - V_2\| < r/\|Y\|$ (which is well defined, thanks to Proposition 3.4.3 on page 61). Then

$$\begin{aligned} \|YV_1 - YV_2\| &= \|Y(V_1 - V_2)\| \\ &\leq \|Y\|\|V_1 - V_2\|, \text{ by the definition of } \|Y\| \\ &< \|Y\|\frac{r}{\|Y\|} = r \end{aligned}$$

So there is an open ball of radius $r/\|Y\|$ around V_1 contained in $Y^{-1}U$. That is, $Y^{-1}U$ is open. Thus Y is a continuous map. $\square$

3.5 FIRST MAIN THEOREM

3.5.1 STATEMENT AND PROOF

We're almost ready to state and prove our first main theorem. Before we do, though, we'll need another of the fundamental results from Functional Analysis, the Hahn-Banach Theorem. Morrison, in his well-written book on Functional Analysis (2001), works through various versions of the Hahn-Banach Theorem. This is the extension form, which allows us to extend a linear functional on a subspace of $C(\mathcal{E})$ to all of $C(\mathcal{E})$.

Theorem 3.5.1 (Hahn-Banach Theorem, Extension Form). *Suppose X is a real linear space and S is a linear subspace of X and $p: X \rightarrow \mathbb{R}$ is*

Subadditive (that is, $p(x_1 + x_2) \leq p(x_1) + p(x_2)$ for all $x_1, x_2 \in X$)

Nonnegatively subhomogeneous (i.e., $p(\lambda x) \leq \lambda p(x)$ for all $\lambda \geq 0$ and $x \in X$).

Let $f: S \rightarrow \mathbb{R}$ be a linear functional that satisfies $f(s) \leq p(s)$ for all $s \in S$. Then f may be extended in a linear fashion to a functional F defined on all of X in such a manner that the extension satisfies $F(x) \leq p(x)$ for all $x \in X$.

Proof. See [Morrison \(2001, p. 65\)](#). □

Theorem 3.5.2 (First Main Theorem). *The tensegrity $G(p)$ is partially $\mathcal{X}$ -bar equivalent if and only if $G(p)$ has a semipositive stress.*

Proof. Let us begin by assuming that $G(p)$ has a semipositive stress, μ . We wish to show that $G(p)$ has no strictly positive motion. Let $V \in \mathcal{X}$ be such that $YV \in C(\mathcal{E})^+$. Since μ is a stress, we know that

$$0 = \mu(YV) = \int_{\mathcal{E}} YV \, d\mu = \int_{\text{supp } \mu} YV \, d\mu.$$

So we must have $YV = 0$ almost everywhere on $\text{supp } \mu$, a set which, by [Lemma 3.3.11](#) on page [57](#), is nonempty due to μ being semipositive.

Conversely, let us assume that $G(p)$ is partially $\mathcal{X}$ -bar equivalent. Then $Y(\mathcal{X}) \cap \text{int } C(\mathcal{E})^+ = \emptyset$ (by definition). So $Y(\mathcal{X})$ contains no strictly positive motions, that is, for any motion $V \in \mathcal{X}$, there must be some $e_V \in \mathcal{E}$ such that $YV(e_V) = 0$. We also note that since $\mathcal{X}$ is a subspace, $Y(\mathcal{X})$ must be a subspace of $C(\mathcal{E})$, so $Y(\mathcal{X}) \cap \text{int } C(\mathcal{E})^- = -(-Y(\mathcal{X}) \cap \text{int } C(\mathcal{E})^+) = -(Y(\mathcal{X}) \cap \text{int } C(\mathcal{E})^+) = \emptyset$.

Let $\mathfrak{s}: C(\mathcal{E}) \rightarrow \mathbb{R}$ by $\mathfrak{s}(f) = \sup f$ for all $f \in C(\mathcal{E})$. We note that if f_1 and f_2 are in $C(\mathcal{E})$ and $a \in (0, \infty)$, then

$$\mathfrak{s}(f_1 + f_2) = \sup_{e \in \mathcal{E}} \{(f_1 + f_2)(e)\} \leq \sup_{e \in \mathcal{E}} \{f_1(e)\} + \sup_{e \in \mathcal{E}} \{f_2(e)\} = \mathfrak{s}(f_1) + \mathfrak{s}(f_2),$$

and

$$\mathfrak{s}(af_1) = \sup_{e \in \mathcal{E}} \{af_1(e)\} = a \sup_{e \in \mathcal{E}} \{f_1(e)\} = a\mathfrak{s}(f_1),$$

so $\mathfrak{s}$ is subadditive and nonnegatively subhomogeneous.

Now, by Lemma 3.3.2 on page 53, $\mathbf{1} \in \text{int } C(\mathcal{E})^+$, so, by hypothesis, $\mathbf{1} \notin Y(\mathcal{X})$. We can define $\hat{\mu}: \text{span}\{Y(\mathcal{X}), \mathbf{1}\} \rightarrow \mathbb{R}$ by $\hat{\mu}(YV + \alpha\mathbf{1}) = \alpha$. What can we say about $\hat{\mu}$? First, it is clearly linear. Secondly, it is certainly zero on all of $Y(\mathcal{X})$. How does it compare to $\mathfrak{s}$? Well, we know that any element YV of $Y(\mathcal{X})$ must be zero somewhere (so as to avoid being in either $\text{int } C(\mathcal{E})^+$ or $\text{int } C(\mathcal{E})^-$), so at that point $YV + \alpha\mathbf{1}$ equals α . That gives us that $\mathfrak{s}(YV + \alpha\mathbf{1}) \geq \alpha = \hat{\mu}(YV + \alpha\mathbf{1})$. Hence $\hat{\mu}$ is dominated by $\mathfrak{s}$.

By the Hahn-Banach theorem (Theorem 3.5.1 on page 64), then, we can extend $\hat{\mu}$ to some linear functional μ on all of $C(\mathcal{E})$ that is also dominated by $\mathfrak{s}$.

Let's calculate

$$\|\mu\| = \sup_{\|f\|=1} |\mu(f)|$$

Remembering that $\|f\| = \sup_{e \in \mathcal{E}} |f(e)|$, we see that if $\|f\| = 1$, we must have $\mathfrak{s}(f) \leq 1$ and $\mathfrak{s}(-f) \leq 1$. Since $\mathfrak{s}$ dominates μ , then, we must have $\mu(f) \leq \mathfrak{s}(f) \leq 1$ and $\mu(-f) \leq \mathfrak{s}(-f) \leq 1$, giving us (by the linearity of μ) $\mu(f) = -\mu(-f) \geq -1$. So $\|\mu\| \leq 1$. Since μ is linear and bounded, by the Riesz Representation Theorem (Theorem 3.3.5 on page 55), μ is a regular measure on $\mathcal{E}$. We need to show that $\mu(f) \geq 0$ for all $f \in C(\mathcal{E})^+$ and that $\mu \neq \mathbf{0}$.

Suppose that $f \in C(\mathcal{E})^+$. Then $0 \geq \mathfrak{s}(-f) \geq \mu(-f)$. But $\mu(-f) \leq 0$ implies that $\mu(f) \geq 0$ by the linearity of μ . Since $\mu(\mathbf{1}) = 1$, μ is nonzero, and since $\mu(f) \geq 0$ for all $f \in C(\mathcal{E})^+$, μ is (at least) a semipositive measure. Finally, as $\mu(YV) = 0$ for all $V \in Y(\mathcal{X})$, μ is a semipositive stress and our proof is complete. $\square$

3.5.2 AN EXAMPLE

Perhaps we could benefit from an example. Suppose we have a tensegrity like the one in Figure 3.5.1 on the following page. In the figure, we have a semipositive stress for

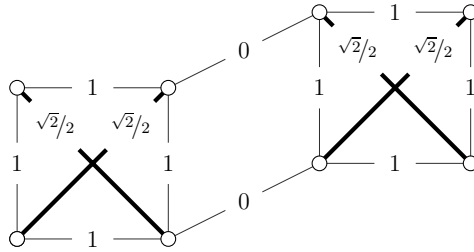


Figure 3.5.1: A tensegrity with a semipositive stress (shown) but no strictly positive one.

this tensegrity, so there is no motion that strictly increases the lengths of all struts and decreases the lengths of all cables.

On the other hand, if we were to put a positive weight on the two edges which are currently marked 0, there would be no way to get a zero vector sum at the vertices they touch. As we saw in Proposition 2.1.17 on page 22, that says there is no strictly positive stress. So the tensegrity is partially $\mathcal{X}$ -bar equivalent (each of the squares is $\mathcal{X}$ -bar equivalent), but it is not $\mathcal{X}$ -bar equivalent as a whole.

Now, Theorem 3.5.2 doesn't seem particularly helpful if there are bars in the tensegrity. After all, one could create a semipositive stress by putting equal positive weights on the strut and cable portions of one of the bars and zeros everywhere else. That is true, but there are a couple of points worth noting. First, there are many continuous tensegrities which do not have bars. Connelly et al. (2003), for example, used just such a theorem on a tensegrity which had no bars.

Secondly, in some cases we can narrow $\mathcal{X}$ down to variations which would not change the bars, and then remove the bars and analyze the tensegrity without them. This has to be done carefully, as removing the bars may destroy the compactness of $\mathcal{E}$ (it would, for example, for the tensegrity in Figure 3.2.6 on page 51), but it is useful in other cases.

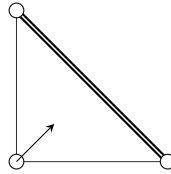


Figure 3.5.2: A tensegrity that has no semipositive stress (and hence has a strictly positive motion) if $\mathcal{X}$ only contains the variations that do not change the length of the bar.

As an example, the tensegrity in Figure 3.5.2 has no motion that strictly shortens all cables and lengthens all struts, because that would require both lengthening and shortening the bar. But if $\mathcal{X}$ excludes those variations that would change the length of the bar, the tensegrity has no semipositive stress. By Theorem 3.5.2, then, there is a strictly positive motion of the other edges.

3.6 SECOND MAIN THEOREM

Our original goal was to move Roth & Whiteley's lemma into the continuous world. We haven't quite accomplished it. The difference between what we have and what we were after is the same as the difference between Gordan's theorem and Stiemke's (shown on page 38). Let's take another pass at it.

3.6.1 ONE DIRECTION

One direction of the theorem is easy to do.

Theorem 3.6.1 (Second Main Theorem). *If a tensegrity has a strictly positive stress, then it is $\mathcal{X}$ -bar equivalent.*

Proof. Let $G(p)$ be a tensegrity. $G(p)$ having a strictly positive stress means that there exists some strictly positive measure $\mu \in Y(\mathcal{X})^\perp$.

Let V be any element of $\mathcal{X}$ such that $YV \in C(\mathcal{E})^+$. Since $\mu \in Y(\mathcal{X})^\perp$, we must have $\mu(YV) = 0$. On the other hand, μ is strictly positive and $YV \in C(\mathcal{E})^+$, so $\mu(YV) = 0$ implies $YV = \mathbf{0}$. $\square$

3.6.2 PARTIALLY $\mathcal{X}$ -BAR EQUIVALENT

Let's pause for a moment to defend the choice of “partially $\mathcal{X}$ -bar equivalent” as a term.

Corollary 3.6.2. *If $G(p)$ is partially $\mathcal{X}$ -bar equivalent, then some subtensegrity of $G(p)$ is $\mathcal{X}$ -bar equivalent.*

Proof. Let $G(p)$ be partially $\mathcal{X}$ -bar equivalent. By Theorem 3.5.2, $G(p)$ has a semi-positive stress μ . Then $\text{supp } \mu$ is nonempty, by Lemma 3.3.11, and closed, by Lemma 3.3.8, (and hence compact as the closed subspace of a compact space, see, for example [Munkres \(2000, p. 165\)](#)) and has positive measure, since $\mu \neq \mathbf{0}$.

So let $S(p)$ be the subtensegrity we get from $G(p)$ by keeping the vertex set $\mathcal{V}$ intact but by restricting the edgeset to $\text{supp } \mu$. Lemma 3.3.10 gives us that μ is a strictly positive measure for $S(p)$, so we need only show that it is a stress. But since μ is a stress for $G(p)$, we already have

$$0 = \mu(YV) = \int_{\text{supp } \mu} YV,$$

so μ is a stress for $S(p)$ as well. Hence, by Theorem 3.6.1, $S(p)$ is $\mathcal{X}$ -bar equivalent. $\square$

We note, in passing, that [Figure 3.2.1](#) on page 42 does not give us a counterexample to this theorem, as $\mathcal{E}$ in that case is not compact.

3.6.3 THE OTHER DIRECTION

We have that having a strictly positive stress implies $\mathcal{X}$ -bar equivalence. Showing that the converse is true for continuous tensegrities is surprisingly difficult. Here's a

little idea of why. Figure 3.6.1 shows a tensegrity which resembles the crossed square of Section 2.2. This one differs, however, in that it has a bar in place of one of the

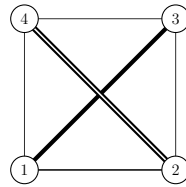


Figure 3.6.1: A non-minimal crossed square.

struts. We already know that the crossed square is bar equivalent, what effect does “extra cable” of the bar have? In Section 2.2, we found that giving a weight of 1 to every edge provided a strictly positive stress. By inspecting $Y^\top$ (on page 25), we can discover that elements of its kernel must have the same weight on every edge, so our μ is, up to scaling, the only strictly positive stress for the crossed square.

By changing that strut to a bar, though, we changed Y . It now has another row corresponding the new “cable”.

$$\hat{Y} = \begin{bmatrix} 1 & 0 & -1 & 0 & 0 & 0 & 0 & 0 \\ -1 & -1 & 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 1 & 0 & -1 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 0 & -1 & 1 \\ 0 & 0 & 0 & 0 & -1 & 0 & 1 & 0 \\ 0 & 0 & -1 & 1 & 0 & 0 & 1 & -1 \end{bmatrix}$$

A little arithmetic and we find that the kernel of $\hat{Y}^\top$ is now a family of stresses which look like $\begin{bmatrix} a & a & a & a & b & a & (b-a) \end{bmatrix}^\top$. There is now an interval’s worth of strictly positive stresses. For any given positive value of b , a can take on any value in the interval $(0, b)$.

In [Subsection 3.8.2](#) on page 89, we'll return to this example and see what that says about the tensegrity, but at the moment we're more interested in what it says about μ .

While the values $a \in (0, b)$ give us strictly positive stresses, the values $a = 0$ and $a = b$ give us semipositive stresses, which lie on the boundary of $C^*(\mathcal{E})^+$. We could turn those semipositive stresses into strictly positive ones by “rotating” the stress into the interior of $C^*(\mathcal{E})$. In the continuous case, [Theorem 3.5.2](#) on page 65 gives us semipositive stresses, so why not do something similar and “rotate” those stresses into the interior of $C^*(\mathcal{E})^+$? The answer is fairly fundamental: In the infinite case, $C^*(\mathcal{E})^+$ has no interior.

Suppose X is any compact Hausdorff space and let $C(X)$, $C^*(X)$ and $C^*(X)^+$ be defined for X just as $C(\mathcal{E})$, $C^*(\mathcal{E})$ and $C^*(\mathcal{E})^+$ are defined for $\mathcal{E}$ for continuous tensegrities. Then,

Theorem 3.6.3. *$C^*(X)^+$ has interior if and only if X is finite.*

Proof. Suppose X is finite. The elements of $C(X)$ and $C^*(X)$ are scalar fields on X . Let $\mu \in C^*(X)^+$ with $\mu(x) > 0$ for all $x \in X$ and let x_0 denote some element of X on which μ takes its smallest value. Now let $\eta \in C^*(X)$ with $\|\eta - \mu\| < \mu(x_0)$, that is,

$$\sup_{\|f\|=1} |(\eta - \mu)(f)| < \mu(x_0).$$

Then certainly $|(\eta - \mu)(f)| < \mu(x_0)$ for all f of norm 1, including those functions which map one element of X to 1 and all of the others to 0. So for every $x \in X$, we have

$$\begin{aligned} |(\eta - \mu)(x)| < \mu(x_0) &\Rightarrow -\mu(x_0) < \eta(x) - \mu(x) < \mu(x_0) \\ &\Rightarrow \mu(x) - \mu(x_0) < \eta(x) < \mu(x) + \mu(x_0) \end{aligned}$$

and since $\mu(x) \geq \mu(x_0)$ for all x , we have $\eta(x) > 0$ for all $x \in X$. Thus in the finite case, $C^*(X)^+$ has nonempty interior — the strictly positive elements of $C^*(X)$ form its interior.

Let's consider the general case. Suppose $\mu \in \text{int } C^*(X)^+$. At the very least that means that $\mu(X) \geq 0$, so by scaling we can assume that $\mu(X) = 1$. Since $C^*(X)$ has the metric topology, μ being in the interior of $C^*(X)^+$ also means that there must be some $\delta > 0$ such that whenever $\eta \in C^*(X)$ with $\|\mu - \eta\| < \delta$, we have $\eta \in C^*(X)^+$.

Now either there exists some $\varepsilon > 0$ such that $\mu(U) \geq \varepsilon$ for every nonempty open $U \subset X$, or not.

Let us assume first that such an ε exists. Then any collection of nonempty, pairwise-disjoint open sets on X must have at most $\lfloor 1/\varepsilon \rfloor$ members (where $\lfloor x \rfloor$ designates the greatest integer less than or equal to x).

Let m be the maximum number of sets in any such collection and let $U_1, \dots, U_m$ be a collection with that maximal number of members. Now the union of the closures of the U_i ,

$$\bigcup_{i=1}^m \overline{U_i}$$

must contain all of X . Otherwise $X \setminus \bigcup \overline{U_i}$ would be another open set disjoint from each of the U_i , contradicting the maximality of m .

There are only finitely many U_i , so if there are infinitely many elements in X , then by the Pigeon-hole Principle (see, for example, [Landau \(1958, p. 78\)](#)), at least one of the $\overline{U_i}$ must contain infinitely many of them. Without loss of generality, we'll take that set to be $\overline{U_m}$.

Let x_1 and x_2 be two distinct elements in $\overline{U_m}$. Since X is Hausdorff, there exist disjoint open sets S_1 and S_2 which separate them. Since x_1 and x_2 are in the closure of U_m , it must be true that $S_1 \cap U_m \neq \emptyset$ and $S_2 \cap U_m \neq \emptyset$. Hence, by removing U_m from the collection and replacing it with $S_1 \cap U_m$ and $S_2 \cap U_m$, we have a new collection of nonempty, pairwise-disjoint open sets with $m + 1$ elements, again contradicting the maximality of m .

Thus, if there exists such an ε , there are only a finite number of elements in X .

On the other hand, suppose that there is no such ε . Then we can select some U_0 such that $\mu(U_0) < \delta/4$ and let μ_δ be a semipositive measure such that $\mu_\delta(U_0) = \delta/2$, but $\mu_\delta(U) = 0$ whenever $U \subset X$ with $U \cap U_0 = \emptyset$ (for example, since U_0 is nonempty, $\mu(U_0) > 0$, so select any $x \in U_0$ and put an atom of size $\delta/2$ on it and zeros everywhere else). We note that since $(X \setminus U_0) \cap U_0 = \emptyset$, $\mu_\delta(X) = \mu_\delta(U_0) = \delta/2$. And, since μ_δ is semipositive,

$$\|\mu_\delta\| = \sup_{\|f\|=1} \mu_\delta(f) = \mu_\delta(\mathbf{1}) = \mu_\delta(U_0) = \delta/2.$$

Now consider $\eta := \mu - \mu_\delta$. By construction,

$$\|\mu - \eta\| = \|\mu - (\mu - \mu_\delta)\| = \|\mu_\delta\| = \delta/2 < \delta,$$

so η lies within a ball of radius δ about μ . But

$$\eta(U_0) = \mu(U_0) - \mu_\delta(U_0) < \delta/4 - \delta/2 = -\delta/4 < 0$$

so $\eta \notin C^*(X)^+$.

Thus $C^*(X)^+$ has interior if and only if X is finite. $\square$

As that route to a strictly positive stress has proved unproductive, let's briefly set aside our quest to show that $\mathcal{X}$ -bar equivalent tensegrities have strictly positive stresses and take a look at the results of what we have accomplished so far.

3.7 TWO EXAMPLES

3.7.1 THE SETUP

We can explore the new theorems with a couple of examples. In the first, shown in [Figure 3.7.1](#) on the following page, the vertices form a unit circle and struts are placed antipodally at all points. Our design variations will be only those variations which do not even locally stretch or shrink the curve of vertices. Since the circle cannot stretch or shrink, this example certainly appears to be $\mathcal{X}$ -infinitesimally rigid and so we would expect it to be $\mathcal{X}$ -bar equivalent.

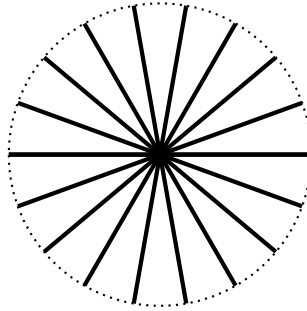


Figure 3.7.1: The circle-of-struts example.

We get the second example by removing a little more than half the struts of the first example (see [Figure 3.7.2](#)). We'll leave those struts that touch the circle from ε to $\pi/2 - \varepsilon$ for some $\varepsilon \in (0, \pi/4)$. That should allow enough freedom to establish a strictly positive motion, even with our requirement about maintaining length locally on the vertex set.

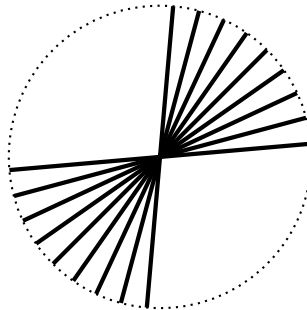


Figure 3.7.2: The almost-half-a-circle-of-struts example.

In both cases, we can parameterize $\mathcal{V}$ and $\mathcal{E}$ by angle (since angle and arclength are the same on a unit circle). The values for $\mathcal{V}$ will lie in $\mathbb{R}/2\pi$ and for $\mathcal{E}$ in $\mathbb{R}/\pi$.

3.7.2 A NEW $\mathcal{X}$

Much of our preparation for this example will center around the set of design variations, $\mathcal{X}$. Specifically, we need to identify those variations which are local isometries of the vertex curve (to remind us that it is a curve, we'll call it γ instead of $\mathcal{V}$).

For a little while, we'll even forget that our examples have circles for vertex curves, and see how general we can let γ and $\mathcal{X}$ be.

We need to be able to define length. To do that, we approximate γ with inscribed polygons or polygonal lines and take its length to be the supremum of the lengths of all such approximations. Here's a more rigorous definition:

Definition 3.7.1 (Strichartz (1995, p. 611)). The *length* of a continuous curve $\gamma: [a, b] \rightarrow \mathbb{R}^n$ is the sup of $\sum_{j=1}^N \|\gamma(t_j) - \gamma(t_{j-1})\|$ taken over all partitions $a = t_0 < t_1 < \dots < t_N = b$ of the interval $[a, b]$, where $\|\gamma(t_j) - \gamma(t_{j-1})\|$ is the distance between the points $\gamma(t_j)$ and $\gamma(t_{j-1})$ in $\mathbb{R}^n$. If the length is finite, we say the curve is *rectifiable*.

We'll require that $p(\gamma)$ be a simple, rectifiable curve with a finite number of connected components and total length $\ell < \infty$.

One advantage of a rectifiable curve is that we can give it a nice parameterization. Using our definition of length, we can define:

Definition 3.7.2. An *arclength parameterization* of a curve $\gamma: [a, b] \rightarrow \mathbb{R}^n$ of length ℓ is a continuous map $t(s): [0, \ell] \rightarrow [a, b]$ such that at any point $s \in [0, \ell]$, the length of γ from $\gamma(t(0)) = \gamma(a)$ to $\gamma(t(s))$ is equal to s .

Proposition 3.7.3. *Every rectifiable curve admits an arclength parameterization.*

Proof. See, for example, Strichartz (1995, p. 616) or Burago et al. (2001, p. 46). $\square$

Now suppose we have some arclength parameterized $\gamma: [0, \ell] \rightarrow \mathbb{R}^n$. What more can we say about γ ? It turns out we can say some strong things about continuity. Here are four different types of continuity, each stronger than the previous:

Definition 3.7.4 (see, for example, Royden (1968, p. 44)). A function $f(x)$ is *continuous* if, for every $\varepsilon > 0$ and every x , there exists a $\delta > 0$ such that whenever $|x - y| < \delta$, $|f(x) - f(y)| < \varepsilon$.

Definition 3.7.5 (see, for example, Royden (1968, p. 135)). A function $f(x)$ is *uniformly continuous* if, for every $\varepsilon > 0$ there exists a $\delta > 0$ such that for every x , whenever $|x - y| < \delta$, $|f(x) - f(y)| < \varepsilon$.

Definition 3.7.6 (see, for example, Folland (1999, p. 105)). A function $f(x)$ is *absolutely continuous* if, for every $\varepsilon > 0$ there exists a $\delta > 0$ such that for every finite collection of pairwise-disjoint intervals (x_i, y_i) with $\sum |x_i - y_i| < \delta$ we have $\sum |f(x_i) - f(y_i)| < \varepsilon$.

Definition 3.7.7 (see, for example, Folland (1999, p. 108) or Conway (1990, p. 25)). A function f is *Lipschitz continuous* (or just *Lipschitz*) if there exists some constant $k \geq 0$ (called the *Lipschitz constant*) such that for all x and y , we have $|f(x) - f(y)| \leq k|x - y|$.

We claimed that each of the types of continuity was stronger than the previous. Here's proof of that relationship for the last two:

Proposition 3.7.8. *If a map f is Lipschitz, it is absolutely continuous.*

Proof. Let $\varepsilon > 0$ be given. Select $\delta = \varepsilon/k$ where k is the Lipschitz constant of f . Then, for every finite collection of intervals, (x_i, y_i) , with $\sum |x_i - y_i| < \delta$, we have that

$$\begin{aligned} \sum |f(x_i) - f(y_i)| &\leq \sum k|x_i - y_i| \text{ since } f \text{ is Lipschitz} \\ &= k \sum |x_i - y_i| < k\delta = \varepsilon. \end{aligned}$$

□

Now the converse of Proposition 3.7.8 is not generally true. Here's a case in which it is not.

Claim 3.7.9. *The function $\sqrt{x}$ is absolutely continuous but not Lipschitz continuous on $[0, 1]$.*

Proof. Consider the function $\sqrt{x}$. Let's show that $\sqrt{x}$ is absolutely continuous.

For any given $\varepsilon > 0$, we can select $\delta = \varepsilon^2$. Now $\sqrt{x}$ has monotonically decreasing derivative on $(0, \infty)$. So the worst-case collection of intervals of total length less than δ will be the single interval $(0, \rho)$ for any $\rho < \delta$. But then $|\sqrt{0} - \sqrt{\rho}| < \sqrt{\delta} = \sqrt{\varepsilon^2} = \varepsilon$.

On the other hand, $\sqrt{x}$ is not Lipschitz continuous, since the expression $\frac{|\sqrt{0} - \sqrt{x}|}{|0 - x|} = \frac{\sqrt{x}}{x} = \frac{1}{\sqrt{x}}$ can be made arbitrarily large by getting x near enough to 0. $\square$

John Sullivan, in his survey paper, “Curves of Finite Total Curvature”, gives the following proposition.

Proposition 3.7.10. *A curve is rectifiable if and only if it admits a Lipschitz parameterization.*

Proof. See [Sullivan \(2006, p. 3\)](#). $\square$

What he is noting, in the “only if” portion of his proof, is that an arclength parameterized curve, while not necessarily differentiable, is Lipschitz (with Lipschitz constant 1). So we have both Lipschitz continuity and absolute continuity of γ simply due to its being rectifiable. What more can we get?

Lemma 3.7.11. *If $f: [a, b] \rightarrow \mathbb{R}^n$ is absolutely continuous, then f is rectifiable.*

Proof. Here we flesh out a proof of Royden ([1968](#), p. 105). Since f is absolutely continuous, there exists some $\delta > 0$ such that for every collection of pairwise-disjoint intervals (x_i, y_i) with $\sum |x_i - y_i| < \delta$, we have $\sum \|f(x_i) - f(y_i)\| < 1$. Let $m = \lfloor \frac{b-a}{\delta} \rfloor$. Given any partition $a < t_0 < t_1 < \dots < t_N = b$, we can split the intervals (t_i, t_{i+1}) (if necessary) and group them in such a way that we have $m + 1$ groups: m whose total length falls in the range $(\delta - \frac{1}{m}\delta, \delta)$ and 1 with total length less than $m\frac{1}{m}\delta = \delta$.

By the triangle inequality, splitting intervals in a partition will never decrease the value of

$$\sum_{j=1}^N \|f(t_j) - f(t_{j-1})\|$$

so we know that for any partition, that sum is at most $(m+1)(1) = m+1$, so the length is also bounded above by $m+1 < \infty$. $\square$

Wait a moment. How can it be true that rectifiable implies Lipschitz implies absolutely continuous implies rectifiable and yet absolutely continuous is not equivalent to Lipschitz?

The answer is that rectifiability is an attribute of the trace of a curve in space, while Lipschitz and absolute continuity are attributes of the function mapping out that trace. An absolutely continuous function maps out a curve which admits a Lipschitz parameterization, but it may not itself be that parameterization. In our example, the trace of the function $x \mapsto \sqrt{x}$ is the interval $[0, \infty)$. And this certainly admits a Lipschitz parameterization, namely the function $x \mapsto x$ (for $x \in [0, \infty)$), which is Lipschitz with Lipschitz constant 1.

What else do we know?

Theorem 3.7.12. *If f is absolutely continuous, then f has a derivative almost everywhere.*

Proof. See, for example, [Royden \(1968, p. 105\)](#). $\square$

So we get derivatives almost everywhere. That allows us to look at the (Lebesgue) integral

$$L(\gamma) = \int_a^b \|\gamma'(s)\| ds.$$

Then we have

Theorem 3.7.13. *The function $L(\gamma)$ gives the arclength of γ (as defined using inscribed polygons) from a to b .*

Proof. See [Burago et al. \(2001, p. 57\)](#). □

We are finally ready to look for those vector fields which preserve length. We need γ to remain rectifiable throughout, so we will restrict ourselves to vector fields which are absolutely continuous.

Let's summarize: γ is a simple, rectifiable curve with a finite number of connected components and total length $\ell < \infty$ and $V: [0, \ell] \rightarrow \mathbb{R}^n$ is an absolutely continuous vector field that moves γ while locally preserving γ 's arclength to first order. What can we say about V ?

Proposition 3.7.14. *If $\gamma(s)$ is a simple, rectifiable, arclength parameterized curve with length ℓ and $V(s)$ is an absolutely continuous variation of $\gamma(s)$ which locally preserves the arclength of $\gamma(s)$ to first order, then $\langle V'(s), \gamma'(s) \rangle = 0$ almost everywhere.*

We'll jump into the proof in a moment, but before we start we'll need a result about Lebesgue integrals.

Theorem 3.7.15 (Dominated Convergence Theorem). *Let $\{f_n\}$ be a sequence of integrable functions such that*

- (a) $f_n \rightarrow f$ almost everywhere and
- (b) *there exists a nonnegative, integrable function g such that $|f_n| \leq g$ almost everywhere for all n .*

Then f is integrable and $\int f = \lim_{n \rightarrow \infty} \int f_n$.

Proof. See, for example, [Folland \(1999, p. 54\)](#) □

Now we're ready for that proof.

Proof of Proposition 3.7.14. For any $a, b \in [0, \ell]$ and $\varepsilon > 0$, we know that the length of the section of $\gamma(s) + \varepsilon V(s)$ between a and b is given by

$$\int_a^b \left\| \frac{d}{ds} (\gamma(s) + \varepsilon V(s)) \right\| ds = \int_a^b \|\gamma'(s) + \varepsilon V'(s)\| ds.$$

If V is to be locally arclength preserving to first order, we must have that

$$\left[\frac{d}{d\varepsilon} \int_a^b \|\gamma'(s) + \varepsilon V'(s)\| ds \right]_{\varepsilon=0} = 0$$

for every choice of $a, b \in [0, \ell]$.

Applying the definition of the derivative, we get

$$\begin{aligned} 0 &= \lim_{\varepsilon \rightarrow 0} \frac{\int_a^b \|\gamma'(s) + \varepsilon V'(s)\| ds - \int_a^b \|\gamma'(s) - (0)V'(s)\| ds}{\varepsilon - 0} \\ &= \lim_{\varepsilon \rightarrow 0} \int_a^b \frac{\|\gamma'(s) + \varepsilon V'(s)\| - \|\gamma'(s)\|}{\varepsilon} ds \end{aligned}$$

Now if we let $\{\varepsilon_n\}$ be a sequence of positive values approaching 0, we see that the functions

$$f_n(s) = \frac{\|\gamma'(s) + \varepsilon_n V'(s)\| - \|\gamma'(s)\|}{\varepsilon_n}$$

approach

$$f(s) = \left[\frac{\partial}{\partial \varepsilon} \|\gamma'(s) + \varepsilon V'(s)\| \right]_{\varepsilon=0}$$

wherever $\gamma'(s)$ and $V'(s)$ are both defined (which is almost everywhere). Now either $\|\gamma'(s) + \varepsilon_n V'(s)\| \geq \|\gamma'(s)\|$ or not. If it is, then

$$\begin{aligned} 0 &\leq \|\gamma'(s) + \varepsilon_n V'(s)\| - \|\gamma'(s)\| \\ &\leq \|\gamma'(s) + \varepsilon_n V'(s) - \gamma'(s)\| = \|\varepsilon_n V'(s)\| = \varepsilon_n \|V'(s)\|. \end{aligned}$$

Otherwise, we have

$$\begin{aligned} 0 &\leq \|\gamma'(s)\| - \|\gamma'(s) + \varepsilon_n V'(s)\| \\ &\leq \|\gamma'(s) - \gamma'(s) - \varepsilon_n V'(s)\| = \|-\varepsilon_n V'(s)\| = \varepsilon_n \|V'(s)\|. \end{aligned}$$

So, by the triangle inequality,

$$\left| \frac{\|\gamma'(s) + \varepsilon_n V'(s)\| - \|\gamma'(s)\|}{\varepsilon_n} \right| \leq \left| \frac{\varepsilon_n \|V'(s)\|}{\varepsilon_n} \right| = \|V'(s)\|$$

for all n . Since $\|V'(s)\|$ is nonnegative and integrable, the Dominated Convergence Theorem tells us that

$$\left[\frac{d}{d\varepsilon} \int_a^b \|\gamma'(s) + \varepsilon V'(s)\| ds \right]_{\varepsilon=0} = \int_a^b \left[\frac{\partial}{\partial \varepsilon} \|\gamma'(s) + \varepsilon V'(s)\| \right]_{\varepsilon=0} ds.$$

Since this is true for *any* choice of a and b , it must be the case that

$$\left[\frac{\partial}{\partial \varepsilon} \|\gamma'(s) + \varepsilon V'(s)\| \right]_{\varepsilon=0} = 0$$

almost everywhere.

Now

$$\begin{aligned} \left[\frac{\partial}{\partial \varepsilon} \|\gamma'(s) + \varepsilon V'(s)\| \right]_{\varepsilon=0} &= \left[\frac{\partial}{\partial \varepsilon} \langle \gamma'(s) + \varepsilon V'(s), \gamma'(s) + \varepsilon V'(s) \rangle^{1/2} \right]_{\varepsilon=0} \\ &= \left[\frac{2 \langle V'(s), \gamma'(s) + \varepsilon V'(s) \rangle}{2 \|\gamma'(s) + \varepsilon V'(s)\|} \right]_{\varepsilon=0} \\ &= \frac{\langle V'(s), \gamma'(s) \rangle}{\|\gamma'(s)\|} \\ &= \langle V'(s), \gamma'(s) \rangle, \end{aligned}$$

so we need $V'(s)$ and $\gamma'(s)$ to be orthogonal almost everywhere. $\square$

Our $\mathcal{X}$, then, is the space of absolutely continuous vector fields V on γ with $V'(s) \perp \gamma'(s)$ for almost all $s \in [0, \ell]$.

We will need to check that $\mathcal{X}$ is a subspace of $\mathbf{VF}(\gamma)$.

Proposition 3.7.16. *$\mathcal{X}$ is a subspace of $\mathbf{VF}(\gamma)$.*

Proof. Let $V, W \in \mathcal{X}$ and $\alpha, \beta \in \mathbb{R}$. Then $\alpha V + \beta W$ is absolutely continuous, and here's why. Let $\varepsilon > 0$. Then (by the absolute continuity of V and W) there exists δ_V and δ_W such that whenever (x_i, y_i) is a finite collection of pairwise-disjoint intervals we have

$$\sum |x_i - y_i| < \delta_V \Rightarrow \sum \|V(x_i) - V(y_i)\| < \frac{\varepsilon}{2|\alpha|}$$

and

$$\sum |x_i - y_i| < \delta_W \Rightarrow \sum \|W(x_i) - W(y_i)\| < \frac{\varepsilon}{2|\beta|}.$$

But then, whenever $\sum |x_i - y_i| < \min\{\delta_V, \delta_W\}$, we have

$$\begin{aligned} \sum \|(\alpha V + \beta W)(x_i) - (\alpha V + \beta W)(y_i)\| \\ &= \sum \|\alpha(V(x_i) - V(y_i)) + \beta(W(x_i) - W(y_i))\| \\ &\leq \sum |\alpha| \|V(x_i) - V(y_i)\| + |\beta| \|W(x_i) - W(y_i)\| \\ &< |\alpha| \frac{\varepsilon}{2|\alpha|} + |\beta| \frac{\varepsilon}{2|\beta|} = \varepsilon. \end{aligned}$$

That gives us absolute continuity. We need to check the orthogonality condition. But

$$(\alpha V + \beta W)'(s) \cdot \gamma'(s) = \alpha V'(s) \cdot \gamma'(s) + \beta W'(s) \cdot V'(s) = 0$$

almost everywhere, so $\alpha V + \beta W \in \mathcal{X}$, and $\mathcal{X}$ is a subspace. $\square$

3.7.3 THE CIRCLE OF STRUTS

Now that we have a grasp on $\mathcal{X}$, let's return to our examples. For both examples, γ' is made up of unit vectors tangent to the circle and pointing counter-clockwise around it. Any vector field $V \in \mathcal{X}$, then, must have V' pointing radially outward or inward on the circle. So everything in $\mathcal{X}$ will be of the nature

$$V(s) = \int_0^s (f(t) \cos t, f(t) \sin t) dt$$

(where $f: \mathbb{R}/\pi \rightarrow \mathbb{R}$ is not necessarily continuous and only need be defined almost everywhere). But why should we believe that every such $V(s)$ is absolutely continuous?

Theorem 3.7.17. *A function f is an indefinite integral if and only if it is absolutely continuous.*

Proof. See, for example, (Royden, 1968, p. 106). $\square$

Since we also need absolute continuity at 0, we'll have the additional requirement that $V(2\pi) = V(0)$. That is,

$$\int_0^{2\pi} (f(t) \cos t, f(t) \sin t) dt = \mathbf{0}.$$

For the circle-of-struts example, let's try building a stress. The simplest option would be to take the measure $d\theta$, which is uniform on $\mathbb{R}/\pi$ (that is, on the struts), and integrates to π . Let's suppose that V is some variation in $\mathcal{X}$ and see what happens.

Now $YV(\theta) = (V(\theta) - V(\theta + \pi)) \cdot (\cos \theta, \sin \theta)$. So we get

$$\begin{aligned} \int_0^\pi YV(\theta) d\theta &= \int_0^\pi (V(\theta) - V(\theta + \pi)) \cdot (\cos \theta, \sin \theta) d\theta \\ &= \int_0^\pi V(\theta) \cdot (\cos \theta, \sin \theta) d\theta + \int_\pi^{2\pi} V(\theta) \cdot (\cos \theta, \sin \theta) d\theta \\ &= \int_0^{2\pi} V(\theta) \cdot (\cos \theta, \sin \theta) d\theta, \end{aligned}$$

which we can integrate by parts to get

$$= V(\theta) \cdot (\sin \theta, -\cos \theta) \Big|_0^{2\pi} - \int_0^{2\pi} V'(\theta) \cdot (\sin \theta, -\cos \theta) d\theta.$$

Since we require that $V(0) = V(2\pi)$, that first term is zero, so we get

$$= \int_0^{2\pi} V'(\theta) \cdot (-\sin \theta, \cos \theta) d\theta.$$

On the other hand, $(-\sin \theta, \cos \theta) = \gamma'(\theta)$ and since $V \in \mathcal{X}$, we know that $V'(\theta) \cdot \gamma'(\theta) = 0$ almost everywhere, so we have $\int V'(\theta) \cdot \gamma'(\theta) = 0$. Hence $d\theta \in Y(\mathcal{X})^\perp$.

So $d\theta$ is a stress. But taking any nonempty, open set of edges and integrating by $d\theta$ will give a positive value, so $d\theta$ is a strictly positive stress. The circle of struts is $\mathcal{X}$ -bar equivalent. If we were to apply Theorem 3.5.2, we would have:

Proposition 3.7.18. *There is no locally arclength preserving motion of a circle that increases all antipodal distances.*

But Theorem 3.6.1 gives us the stronger:

Proposition 3.7.19. *There is no locally arclength preserving motion of a circle that increases any antipodal distance without decreasing some other one.*

3.7.4 ALMOST HALF A CIRCLE OF STRUTS

Now let's consider our other example. In this case, we'd like to show that there is no semipositive stress.

Let's choose a "good" element of $\mathcal{X}$, one in $T(p) \setminus I(p)$. Remember that our tensegrity has less than half the struts that the circle-of-struts example had, so we should be able to expand along those and contract elsewhere, keeping the curve of vertices unstretched. Let's try

$$V_g = \left(\frac{\sin \theta}{2} + \frac{\sin 3\theta}{6}, \frac{\cos \theta}{2} - \frac{\cos 3\theta}{6} \right),$$

which is shown in [Figure 3.7.3](#).

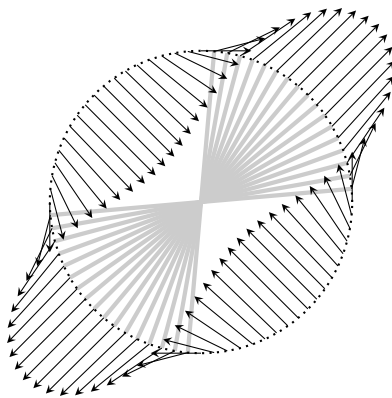


Figure 3.7.3: The almost-half-a-circle-of-struts example with the "good" motion V_g .

Now, YV_g is nonnegative on all of $\mathcal{E}$. In specific,

$$\begin{aligned} p(\theta) - p(\theta + \pi) &= (\cos \theta, \sin \theta) - (\cos(\theta + \pi), \sin(\theta + \pi)) \\ &= (\cos \theta, \sin \theta) - (-\cos \theta, -\sin \theta) \\ &= (2 \cos \theta, 2 \sin \theta) \end{aligned}$$

and similarly,

$$V_g(\theta) - V_g(\theta + \pi) = \left(\sin \theta + \frac{\sin 3\theta}{3}, \cos \theta - \frac{\cos 3\theta}{3} \right),$$

so

$$\begin{aligned}
YV_g(\theta) &= (V_g(\theta) - V_g(\theta + \pi)) \cdot (p(\theta) - p(\theta + \pi)) \\
&= 4 \sin \theta \cos \theta + \frac{2 \sin 3\theta \cos \theta - 2 \cos 3\theta \sin \theta}{3} \\
&= 2 \sin 2\theta + \frac{2 \sin 2\theta}{3} = \frac{8}{3} \sin 2\theta,
\end{aligned}$$

which is strictly positive on the edgeset, that is on $\theta \in [\varepsilon, \pi/2 - \varepsilon]$. But that strict positivity means that for any semipositive μ , we would have μYV_g greater than zero. So there is no semipositive $\mu \in Y(\mathcal{X})^\perp$.

3.7.5 ANY OPEN SET WILL DO

Our success in showing that the almost-half-a-circle-of-struts example is not $\mathcal{X}$ -bar equivalent leads to the next question: If we remove *any* arbitrary open set of edges from the circle of struts, is the result no longer $\mathcal{X}$ -bar equivalent? The answer is yes.

Proposition 3.7.20. *The circle of struts with any nontrivial open set of edges removed is no longer $\mathcal{X}$ -bar equivalent.*

Proof. As before, we consider the set $\mathcal{E}$ to be parameterized by $\theta \in \mathbb{R}/\pi$. Let $U \subset \mathbb{R}/\pi$ represent the open set of edges removed from our example, and let f_U be any continuous function that is positive on $\mathcal{E} = \mathbb{R}/\pi \setminus U$ and for which $\int_0^{2\pi} (f_U(\theta) \cos \theta, f_U(\theta) \sin \theta) d\theta = \mathbf{0}$.

Now define a variation $V_U(\theta) = \int_0^\theta (f_U(t) \cos t, f_U(t) \sin t) dt$. We claim that no semipositive measure, μ , is in $Y(\mathcal{X})^\perp$. After all,

$$\begin{aligned}
YV_U(\theta) &= (f_U(\theta)(\cos \theta, \sin \theta) - f_U(\theta + \pi)(\cos(\theta + \pi), \sin(\theta + \pi))) \cdot \\
&\quad ((\cos \theta, \sin \theta) - (\cos(\theta + \pi), \sin(\theta + \pi))) \\
&= (f_U(\theta) + f_U(\theta + \pi))(\cos \theta, \sin \theta) \cdot 2(\cos \theta, \sin \theta) \\
&= 4f_U(\theta)
\end{aligned}$$

But, by definition, $f_U(\theta) > 0$ on $\mathcal{E}$. So for any semipositive measure μ on $\mathcal{E}$,
 [4] $\mu Y V_U > 0$. □

3.8 MINIMAL $\mathcal{X}$ -BAR EQUIVALENCE

3.8.1 DEFINITION AND EXAMPLE

All we lack to complete our results now is a theorem which says that any $\mathcal{X}$ -bar equivalent tensegrity has a strictly positive measure. We don't have such a theorem for the general case, but there is a class of tensegrities for which we do get the desired result.

Definition 3.8.1. We term a tensegrity $G(p)$ *minimally $\mathcal{X}$ -bar equivalent* if it is $\mathcal{X}$ -bar equivalent but no nonempty subtensegrity formed by removing an open set of edges from it is $\mathcal{X}$ -bar equivalent.

Lemma 3.8.2. *A minimally $\mathcal{X}$ -bar equivalent tensegrity $G(p)$ admits no semipositive stress that is not also strictly positive.*

Proof. By Corollary 3.6.2 on page 69, if the tensegrity admits a semipositive (but not strictly positive) stress μ , then the subtensegrity whose edgeset is $\text{supp } \mu$ is $\mathcal{X}$ -bar equivalent, and $\mathcal{E} - \text{supp } \mu$ is nonempty and open, so $G(p)$ is not minimally $\mathcal{X}$ -bar equivalent. □

Theorem 3.8.3 (Third Main Theorem). *If $G(p)$ is minimally $\mathcal{X}$ -bar equivalent, $G(p)$ has a strictly positive stress.*

Proof. By Theorem 3.5.2, $G(p)$ has a semipositive stress, and by Lemma 3.8.2, that stress must be strictly positive. □

Having seen that a minimally $\mathcal{X}$ -bar equivalent tensegrity cannot admit a (solely) semipositive stress, it seems tempting to conjecture:

Conjecture 3.8.4. *A tensegrity is minimally $\mathcal{X}$ -bar equivalent if and only if it admits only one semipositive stress, up to scaling. Furthermore, that stress is strictly positive.*

What sort of tensegrity is minimally $\mathcal{X}$ -bar equivalent? [Subsection 3.7.5](#) on page 85 showed that the circle of struts is.

Interestingly enough, if we make the circle into a square filled with either antipodal struts or with struts parallel to the edges (the left and right pictures in [Figure 3.8.1](#)), the tensegrity ceases to be $\mathcal{X}$ -bar equivalent ($\mathcal{X}$ is still the motions that are local isometries on γ).

The center picture of [Figure 3.8.1](#) shows a motion that induces a semipositive load on $\mathcal{E}$. In the first case, that load is zero only on the diagonals and in the second case it is zero only on the outermost struts. So any strictly positive stress μ will give

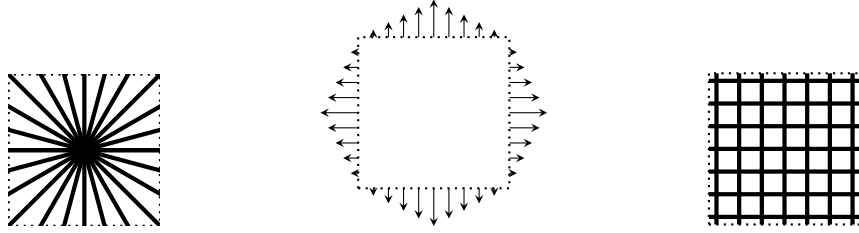


Figure 3.8.1: Two squares of struts and a semipositive motion that works for either.

$\mu YV > 0$ for this motion. On the other hand, the semipositive stress that just assigns an atom to each diagonal (resp. each outermost strut) annihilates $Y(\mathcal{X})$, as we will show below.

Now $\mathcal{X}$ contains those vector fields for which $V'(v) \cdot \gamma'(v) = 0$ except on a set of measure zero. So along the top and bottom lines of vertices, the horizontal component of $V(v)$ cannot vary, while along the sides, its vertical component cannot vary. That means that

$$\begin{aligned} (V((1,0)) - V((0,0))) \cdot (1,0) &= 0 & (V((0,1)) - V((0,0))) \cdot (0,1) &= 0 \\ (V((1,1)) - V((0,1))) \cdot (1,0) &= 0 & (V((1,1)) - V((1,0))) \cdot (0,1) &= 0 \end{aligned}$$

So taking μ to be an atom of size 1 on each outermost strut of the right hand tensegrity, we get

$$\begin{aligned}\mu YV &= (V((1,0)) - V((0,0))) \cdot (1,0) + (V((1,1)) - V((1,0))) \cdot (0,1) + \\ &\quad (V((0,1)) - V((0,0))) \cdot (0,1) + (V((1,1)) - V((0,1))) \cdot (1,0) = 0.\end{aligned}$$

On the other hand, taking μ to give an atom of size 1 on each diagonal strut of the left hand tensegrity, we get

$$\begin{aligned}\mu YV &= (V((1,1)) - V((0,0))) \cdot (1,1) + (V((1,0)) - V((0,1))) \cdot (1,-1) \\ &= (V((1,1)) - V((1,0)) + V((1,0)) - V((0,0))) \cdot (1,1) \\ &\quad + (V((1,0)) - V((1,1)) + V((1,1)) - V((0,1))) \cdot (1,-1) \\ &= (V((1,1)) - V((1,0))) \cdot (1,0) + (V((1,0)) - V((0,0))) \cdot (0,1) \\ &\quad + (V((1,0)) - V((1,1))) \cdot (1,0) + (V((1,1)) - V((0,1))) \cdot (0,-1) \\ &= (V((1,0)) - V((0,0))) \cdot (0,1) + (V((1,1)) - V((0,1))) \cdot (0,-1) \\ &= (V((1,0)) - V((0,0)) + V((0,1)) - V((1,1))) \cdot (0,1) \\ &= (V((1,0)) - V((1,1))) \cdot (0,1) + (V((0,1)) - V((0,0))) \cdot (0,1) = 0\end{aligned}$$

So in each case we have a semipositive stress and the two are partially $\mathcal{X}$ -bar equivalent. Alternatively, that semipositive stress is a strictly positive stress on $\text{supp } \mu$, so the subtensegrities (shown in [Figure 3.8.2](#)) consisting of γ and the mentioned edges are $\mathcal{X}$ -bar equivalent.



Figure 3.8.2: The $\mathcal{X}$ -bar equivalent subtensegrities of the squares of struts.

3.8.2 FINITE IS NOT MINIMAL

It is worth a moment to remember, prior to jumping into the next section, that not all finite tensegrities are minimal. The crossed square with one of the struts replaced with a bar (Figure 3.6.1 on page 70) is just such an example. We saw, in Subsection 3.6.3, that the stresses of that tensegrity look like $\begin{bmatrix} a & a & a & a & b & a & (b-a) \end{bmatrix}^\top$ where $a \in [0, b]$.

If we set $a = 0$, we get a stress which is strictly positive on the bar but zero everywhere else. If we set $a = b$, we get a stress which is strictly positive on the usual crossed square and zero on that new “cable” (see Figure 3.8.3). But these two subtensegrities are the only minimally bar equivalent subtensegrities in this example.⁵

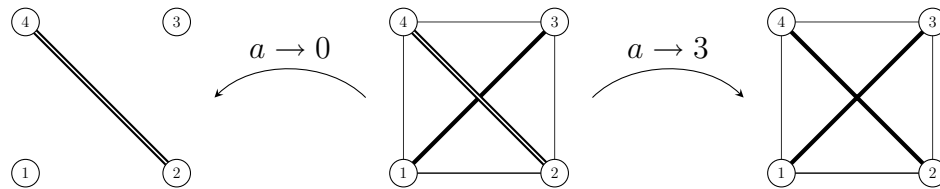


Figure 3.8.3: The non-minimal crossed square with its minimal subtensegrities.

We note that while the tensegrity itself is not minimally bar equivalent, every edge in it belongs to a minimally bar equivalent subtensegrity. Perhaps that is key to having a strictly positive stress.

3.9 COVERED TENSEGRITIES

Let’s consider tensegrities which are covered by minimal subtensegrities. Better yet, let’s consider tensegrities which are almost covered by them.

⁵We can see this by looking at the vertices. If any edge is removed from vertex 1 or vertex 3, all three edges must be removed. That leaves the bar to deal with. If we remove the strut portion of the bar, we must remove all of the other edges at 4 and 2 (and hence also at 1 and 3).

Definition 3.9.1. We say that a tensegrity is *countably covered* by some class of subtensegrities, if there exists a countable, dense subset $\{e_n\} \subset \mathcal{E}$ such that every e_n belongs to a subtensegrity of that class.

Having a countable, dense subset is not an unreasonable thing to ask, as $\mathcal{E}$ is compact and metrizable and hence second countable.⁶

Theorem 3.9.2 (Fourth Main Theorem). *If $G(p)$ is countably covered by subtensegrities that have a strictly positive stress, then $G(p)$ has a strictly positive stress.*

Proof. Let $\{e_n\}$ be the countable, dense subset of edges from the definition of “countably covered”. For each e_n select a subtensegrity G_n with edgeset $\mathcal{E}_n$ of which e_n is a member and which has a strictly positive stress (in the case where e_n belongs to infinitely many subtensegrities, this may involve the Axiom of Choice).

Scale that stress so that its norm is $\frac{1}{2^n}$ and call it $\hat{\mu}_n$. Finally, extend $\hat{\mu}_n$ to be a stress on the entire $G(p)$ by setting, for each open $U \subset \mathcal{E}$, $\mu_n(U) = \hat{\mu}_n(U \cap \mathcal{E}_n)$.

Now we can construct

$$\mu = \sum_{n=1}^{\infty} \mu_n.$$

What do we know about μ ? Well, if $V \in \mathcal{X}$, then every $\mu_n(YV) = 0$ since μ_n is 0 outside $\mathcal{E}_n$ and takes YV to zero on $\mathcal{E}_n$, so $\mu(YV) = 0$. That makes μ a stress. Furthermore, if $U \subset \mathcal{E}$ is open and nonempty, then it must contain at least one of the e_n . But then, by construction, $\mu_n(U) > 0$ and so $\mu(U) > 0$. $\square$

So if we have a tensegrity which is countably covered either by minimally $\mathcal{X}$ -bar equivalent subtensegrities or by finite, $\mathcal{X}$ -bar equivalent subtensegrities (just use the stress from [Roth and Whiteley \(1981\)](#) to build an atomic measure) or by some

⁶This is an exercise (with hint) in [Munkres \(2000, p. 194\)](#). If X is metrizable and compact, then for each positive integer n , we can cover X with balls of size $1/n$. Since X is compact, each such cover has a finite subcover, $\mathcal{A}_n$. The union of all of the $\mathcal{A}_n$ forms a countable basis for the (metric) topology on X and the centers of the balls give a countable, dense subset of X .

combination of the two, then our tensegrity has a strictly positive stress and thus is $\mathcal{X}$ -bar equivalent.

What might such a creature look like? One simple example would be to move the circle-of-struts example “up a dimension”. A “cylinder of struts” is represented in Figure 3.9.1. As with the circle of struts, our design variations are those which are

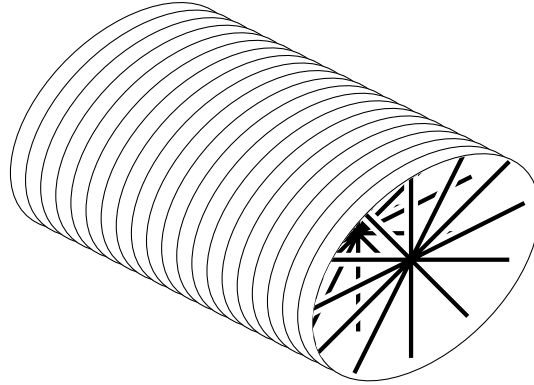


Figure 3.9.1: A cylinder filled with antipodal struts.

local isometries on the surface of the cylinder.

Of course, any given strut in the cylinder of struts is part of a circle of struts. So every strut is contained in a minimally $\mathcal{X}$ -bar equivalent subtensegrity and our tensegrity is (more than) countably covered in minimally $\mathcal{X}$ -bar equivalent subtensegrities. By Theorem 3.9.2, the tensegrity is $\mathcal{X}$ -bar equivalent.

Chapter 4 contains a number of tensegrities which are countably covered by finite subtensegrities.

Now, the “countably covered tensegrities” are a subset of all the $\mathcal{X}$ -bar equivalent tensegrities, but how large a subset is it? That’s not a settled question, but it may be all of them:

Conjecture 3.9.3. *Every $\mathcal{X}$ -bar equivalent continuous tensegrity is countably covered by minimally $\mathcal{X}$ -bar equivalent subtensegrities.*⁷

⁷The matching conjecture, that every $\mathcal{X}$ -bar equivalent continuous tensegrity is countably covered by finite, $\mathcal{X}$ -bar equivalent subtensegrities, is false. One counterexample is

Of course, if that turns out not to be true, we can retreat to the weaker claim:

Conjecture 3.9.4. *Every $\mathcal{X}$ -bar equivalent tensegrity has a strictly positive stress.*

the “on a circle” example of [Section 4.2](#) on page [97](#) when the cable ends are an irrational fraction of the circle circumference apart.

CHAPTER 4

EXAMPLES

Here we have a collection of examples that shed light on various aspects of the subject. Unless stated otherwise, $\mathcal{X} = \text{VF}(\mathcal{V})$ for these examples.

4.1 RECTANGLE

We start with an example that is simple enough to analyze pretty thoroughly. As we are primarily interested in bar equivalence in this paper, we usually give only passing attention to infinitesimal rigidity. But this example is infinitesimally rigid and we take the time to prove it.

The rectangle example lies in $\mathbb{R}^2$ and its has vertices arranged along two intervals, say $(0, 0)$ to $(2, 0)$ and $(0, 1)$ to $(2, 1)$. Each vertex is connected to one directly opposite it by a strut. Every vertex is also connected to the two opposite endpoints by cables (that means that the corner vertices are connected vertically by strut-cable pairs, which is to say, bars). Finally, the two upper corner vertices and the two lower corner vertices are connected by horizontal struts (see [Figure 4.1.1](#) on the next page). Remember that the horizontal struts only connect the corner vertices. They pass through, but do not affect, the other vertices.

We can think of the possible vector fields as four continuous functions from $[0, 2]$ to $\mathbb{R}$, one for the vertical direction and one for the horizontal on each interval. Let's call them f_{vt} (vertical top), f_{vb} , f_{ht} and f_{hb} . Then $\text{im } Y$ is as shown in [Table 4.1](#) (where we've omitted scaling factors since we will only be concerned with signs).

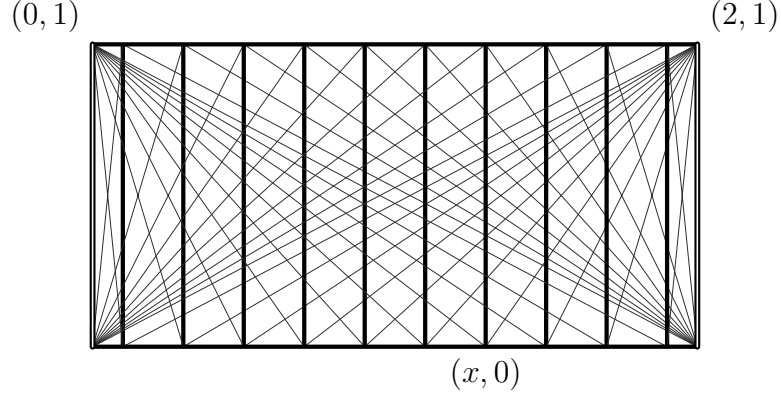


Figure 4.1.1: The rectangle example.

Table 4.1: The values of YV for the different edges of the rectangle tensegrity.

Edge	Value
$(x, 0) \rightarrow (x, 1)$ (strut)	$f_{vt}(x) - f_{vb}(x)$
$(x, 0) \rightarrow (0, 1)$ (cable)	$-x(f_{hb}(x) - f_{ht}(0)) + f_{vb}(x) - f_{vt}(0)$
$(x, 0) \rightarrow (2, 1)$ (cable)	$(2 - x)(f_{hb}(x) - f_{ht}(2)) + f_{vb}(x) - f_{vt}(2)$
$(x, 1) \rightarrow (0, 0)$ (cable)	$-x(f_{ht}(x) - f_{hb}(0)) - f_{vt}(x) + f_{vb}(0)$
$(x, 1) \rightarrow (2, 0)$ (cable)	$(2 - x)(f_{ht}(x) - f_{hb}(2)) - f_{vt}(x) + f_{vb}(2)$
$(0, 0) \rightarrow (2, 0)$ (strut)	$f_{hb}(0) - f_{hb}(2)$
$(0, 1) \rightarrow (2, 1)$ (strut)	$f_{ht}(0) - f_{ht}(2)$

So what are the implications? We'll first examine the associated bar framework. As we work through the families of constraints above, we'll call them struts and cables, but in every case we'll be treating them as equalities, as if they were bars.

Because of the vertical struts, we must have $f_{vt}(x) = f_{vb}(x)$. We can eliminate the Euclidean motions by requiring that $f_{vb}(0) = f_{hb}(0) = f_{ht}(0) = 0$. That results in having $f_{vt}(0) = 0$, since $f_{vt}(0) = f_{vb}(0)$. Our first family of cables then gives us that

$$f_{vb}(x) = x f_{hb}(x) \quad (4.1)$$

and the third set means that

$$f_{vt}(x) = -xf_{ht}(x). \quad (4.2)$$

Since we have changed all edges to bars, the final two struts give us $f_{hb}(2) = f_{ht}(2) = 0$. Now using the second and fourth cable families with $x = 0$, we get

$$2(f_{hb}(0) - f_{ht}(2)) + f_{vb}(0) - f_{vt}(2) \Rightarrow f_{vt}(2) = 0$$

and also $f_{vb}(2) = 0$. For other values of x these give

$$f_{vb}(x) = -(2-x)f_{hb}(x) \quad \text{and} \quad f_{vt}(x) = (2-x)f_{ht}(x). \quad (4.3)$$

So if we combine Equation (4.1) through Equation (4.3), for all values of x we get

$$-xf_{ht}(x) = (2-x)f_{ht}(x) \Rightarrow 2f_{ht}(x) = 0 \Rightarrow f_{ht}(x) = 0$$

and similarly $f_{hb}(x) = 0$ and thus immediately that $f_{vt}(x) = f_{vb}(x) = 0$.

So the bar framework is infinitesimally rigid. Its only motions are the Euclidean motions of $\mathbb{R}^2$. If we can exhibit a strictly positive stress, we'll have that the tensegrity is bar equivalent and hence also infinitesimally rigid. Define μ as follows.

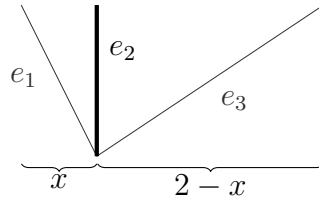


Figure 4.1.2: The edges connecting to an internal vertex in the rectangle.

At any internal point, there are only three edges touching, so we must have something that looks like Figure 4.1.2. Using e_1 , e_2 and e_3 as shown in the figure, we get

$$\mu(e_2) = \frac{\mu(e_1)}{\sqrt{1+x^2}} + \frac{\mu(e_3)}{\sqrt{1+(2-x)^2}} \quad \text{from the vertical stresses, and} \quad (4.4)$$

$$\frac{x}{\sqrt{1+x^2}}\mu(e_1) = \frac{(2-x)}{\sqrt{1+(2-x)^2}}\mu(e_3) \quad \text{from the horizontal ones.} \quad (4.5)$$

If we let $m(e)$ be Lebesgue measure on the edges, inherited from Lebesgue measure on the vertices, and decide that $\mu(e_2) = m(e_2)$, we get

$$\mu(e_1) = \sqrt{1+x^2} \left(m(e_2) - \frac{\mu(e_3)}{\sqrt{1+(2-x)^2}} \right) \quad \text{from Equation (4.4)}$$

and

$$\mu(e_1) = \frac{(2-x)\sqrt{1+x^2}}{x\sqrt{1+(2-x)^2}} \mu(e_3) \quad \text{from Equation (4.5).}$$

Setting these expressions equal results in

$$\frac{(2-x)\mu(e_3)}{x\sqrt{1+(2-x)^2}} = m(e_2) - \frac{\mu(e_3)}{\sqrt{1+(2-x)^2}} \Rightarrow \mu(e_3) = \frac{x\sqrt{1+(2-x)^2}}{2} m(e_2)$$

and directly that

$$\mu(e_1) = \frac{(2-x)\sqrt{1+x^2}}{2} m(e_2).$$

That takes care of all but the endpoints. There are infinitely many cables ending at each corner, but we can integrate to find out how much force is being exerted on those corners.

The forces on the upper left-hand corner all come from those e_1 's above. The e_1 forces resolve into horizontal

$$\frac{x}{\sqrt{1-x^2}} \mu(e_1) = \frac{(2-x)x}{2} m(e_2)$$

and vertical

$$\frac{1}{\sqrt{1-x^2}} \mu(e_1) = \frac{2-x}{2} m(e_2).$$

By symmetry, these are the horizontal and vertical forces at every corner.

Integrating these gives us

$$\int_0^2 \frac{x(2-x)}{2} dm$$

for a total horizontal force of $2/3$ at each corner and

$$\int_0^2 \frac{2-x}{2} dm$$

for a total vertical force of 1 at each corner. So to balance things out, there is an atom of size $4/3$ on the two horizontal struts and an atom of size 2 on the two bars.

We have a strictly positive stress and the rectangle is infinitesimally rigid. Alternatively, we could state it thus:

Proposition 4.1.1. *Any (non-Euclidean) motion of two parallel line segments does at least one of the following:*

1. *Moves two corresponding points on the segments closer together,*
2. *Shortens one or both of the segments, or*
3. *Moves a point on one of the segments farther from one of the endpoints of the opposite segment.*

4.2 ON A CIRCLE

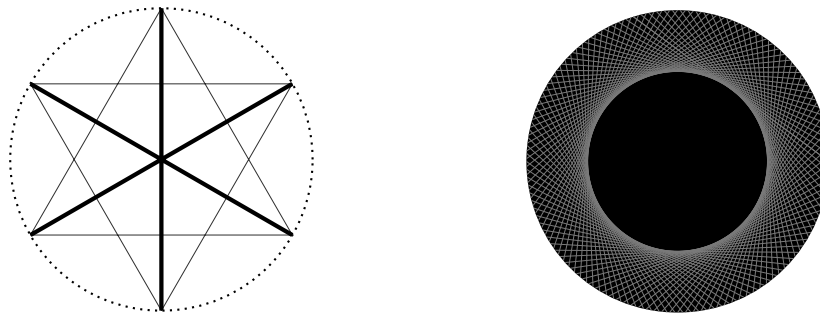


Figure 4.2.1: Two examples of the on-a-circle example.

We next consider a family of examples. Each of these has a unit circle for its vertex set with antipodal points connected by struts. Cables are connected from every point to the two points h units in arc length around the circle from it. If the cable connection distance is a rational fraction of the circumference, the tensegrity consists of infinitely many unconnected finite subtensegrities. If the distance is irrational, it is a single connected tensegrity (see [Figure 4.2.1](#)).

We let m be Lebesgue measure on the edges. Then we create a measure μ on the edges by letting $\mu = m$ for all struts and $\mu = \alpha m$ for all cables. Let's look at the angles involved (see [Figure 4.2.2](#)). If the cables stretch h units around the circle

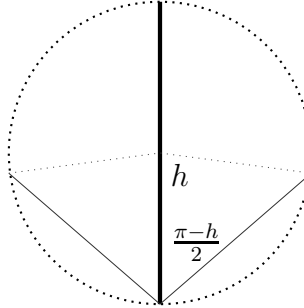


Figure 4.2.2: The angles involved in the on-a-circle example.

(where $0 < h < \pi$), then the cables make an angle of $\frac{\pi-h}{2}$ with the strut meeting at the same vertex.

We want equilibrium at every vertex. The situation is symmetrical across the strut, so we need only check the sum along the strut. If we set $\alpha = \frac{1}{\cos \frac{\pi-h}{2}}$, we get

$$2 dm - 2\alpha \left(\cos \frac{\pi-h}{2} \right) dm = 0 dm$$

at every vertex.

So we have

$$\mu(YV) = \int_{\mathcal{E}} YV d\mu = \int_{v \in \mathcal{V}} V(v) \cdot (0, 0) dm = 0.$$

Thus μ is a strictly positive stress for the tensegrity. Every one of these tensegrities is bar equivalent.

It is worthy of note that if $h \rightarrow 0$, $\alpha \rightarrow \infty$, so we cannot turn this into the example from [Section 3.7](#) on page 73 by a simple limiting process.

4.3 ON A SPHERE

Let's move that example up into three dimensions. We start by examining an octahedral tensegrity that has three struts and twelve cables (see [Section 4.3](#)).

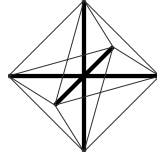


Figure 4.3.1: The octahedron tensegrity.

If we place a compression of unit size on each strut and a tension of $\sqrt{2}/2$ on each cable, then the net force at vertex is zero. So this octahedron is bar equivalent. Now, imagine this octahedron embedded within a sphere as represented in [Figure 4.3.2](#). If we connect every point on a sphere with its antipode by a strut, and connect

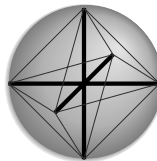


Figure 4.3.2: Octahedra cover the spherical tensegrity.

each point to every point on its associated equator by a cable, we'll have a three-dimensional tensegrity that is covered by octahedra as above and which is thus bar equivalent.

4.4 ON A GENERAL CURVE

Both of the previous two examples have had a lot of symmetry on which we could depend. What if we were to consider something less symmetric? Let's take as our vertex set any simple closed curve in $\mathbb{R}^2$ and perform the same construction we did in [Section 4.2](#). It will be easier, though to think of the cable skip (h in that section) not as a length, but rather as a fraction of the arclength of the vertex curve (see [Figure 4.4.1](#) for two examples with a skip of $1/5$ of the total arclength).

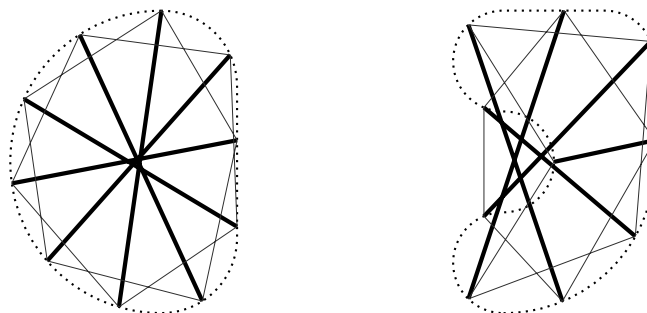


Figure 4.4.1: The on-a-circle example with more general curves.

It seems complicated. Perhaps we should set our sights a little lower at first and see what we can learn.

4.5 UNDERSTANDING A QUADRILATERAL

Suppose we have some quadrilateral that has struts for the diagonals and cables for the edges. In the calculations that follow, we'll use the symbol $\vec{e}_{ij}$ to mean the unit vector pointing from vertex i toward vertex j (so $\vec{e}_{ij} = -\vec{e}_{ji}$) and w_{ij} to mean the value of the stress on the edge connecting vertices i and j . What do we know?

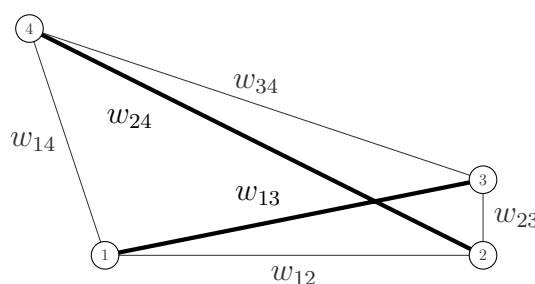


Figure 4.5.1: A general quadrilateral tensegrity.

For one thing, we know that at any vertex, the weighted sum of the cable vectors must equal the weighted strut vector. So the negative of the strut vector must lie in the cone generated by the cable vectors. That is, the angle between the cables must

be less than π on the side where the strut lies and the strut must lie between them. So the quadrilateral must be convex.

Putting the above equalities into symbols, we get

$$\begin{aligned} w_{12}\vec{e}_{12} + w_{14}\vec{e}_{14} + w_{13}\vec{e}_{31} &= 0 & w_{12}\vec{e}_{21} + w_{23}\vec{e}_{23} + w_{24}\vec{e}_{42} &= 0 \\ w_{23}\vec{e}_{32} + w_{34}\vec{e}_{34} + w_{13}\vec{e}_{13} &= 0 & w_{14}\vec{e}_{41} + w_{34}\vec{e}_{43} + w_{24}\vec{e}_{24} &= 0 \end{aligned}$$

Let's change the signs so that the subscripts are in ascending order:

$$w_{12}\vec{e}_{12} + w_{14}\vec{e}_{14} - w_{13}\vec{e}_{13} = 0 \quad (4.6)$$

$$-w_{12}\vec{e}_{12} + w_{23}\vec{e}_{23} - w_{24}\vec{e}_{24} = 0 \quad (4.7)$$

$$-w_{23}\vec{e}_{23} + w_{34}\vec{e}_{34} + w_{13}\vec{e}_{13} = 0 \quad (4.8)$$

$$-w_{14}\vec{e}_{14} - w_{34}\vec{e}_{34} + w_{24}\vec{e}_{24} = 0 \quad (4.9)$$

Since $\vec{e}_{13}$ must lie between $\vec{e}_{12}$ and $\vec{e}_{14}$, it splits the plane into two half-planes, each containing one of those vectors. There are two candidates for “the vector normal to $\vec{e}_{13}$ ”. We'll choose the one that lies in the same half-plane as $\vec{e}_{14}$.

Equation (4.6) can be thought of as two equations in three variables:

$$w_{12}\vec{e}_{12} \cdot \vec{e}_{13} + w_{14}\vec{e}_{14} \cdot \vec{e}_{13} - w_{13} = 0 \quad (4.10)$$

$$w_{12}\vec{e}_{12} \cdot \vec{e}_{13}^\perp + w_{14}\vec{e}_{14} \cdot \vec{e}_{13}^\perp = 0. \quad (4.11)$$

We note that the two dot products in Equation (4.10) are both positive, while our choice of $\vec{e}_{13}^\perp$ makes $\vec{e}_{14} \cdot \vec{e}_{13}^\perp$ positive and $\vec{e}_{12} \cdot \vec{e}_{13}^\perp$ negative.

We can solve this system of equations:

$$w_{12} = \frac{\vec{e}_{14} \cdot \vec{e}_{13}^\perp}{-\vec{e}_{12} \cdot \vec{e}_{13}^\perp} w_{14} \Rightarrow w_{13} = \left[(\vec{e}_{12} \cdot \vec{e}_{13}) \frac{\vec{e}_{14} \cdot \vec{e}_{13}^\perp}{-\vec{e}_{12} \cdot \vec{e}_{13}^\perp} + \vec{e}_{14} \cdot \vec{e}_{13} \right] w_{14}.$$

Letting θ_{jik} be the angle between $\vec{e}_{ij}$ and $\vec{e}_{ik}$ we can rewrite these in terms of sines and cosines:

$$\begin{aligned} w_{12} &= \frac{\sin \theta_{314}}{\sin \theta_{213}} w_{14} \\ w_{13} &= \left[\cos \theta_{213} \frac{\sin \theta_{314}}{\sin \theta_{213}} + \cos \theta_{314} \right] w_{14} = \frac{\sin \theta_{214}}{\sin \theta_{213}} w_{14} \end{aligned}$$

But since $0 < \theta_{213}, \theta_{314}, \theta_{214} < \pi$ we know that w_{12} , w_{13} and w_{14} will all have the same sign (and can certainly be positive).

Similar calculations on [Equation \(4.8\)](#) give us that

$$w_{23} = \frac{\sin \theta_{134}}{\sin \theta_{132}} w_{34} \quad \text{and} \quad w_{13} = \frac{\sin \theta_{234}}{\sin \theta_{132}} w_{34},$$

again with the same sign. So we have a one-parameter family of solutions:

$$w_{12} = \frac{\sin \theta_{314}}{\sin \theta_{214}} w_{13} \quad w_{14} = \frac{\sin \theta_{213}}{\sin \theta_{214}} w_{13} \quad w_{23} = \frac{\sin \theta_{134}}{\sin \theta_{234}} w_{13} \quad w_{34} = \frac{\sin \theta_{132}}{\sin \theta_{234}} w_{13}$$

There are some points worth noting about this. First, the “double arrow” consisting of a strut and the four cables to which it connects (like the ones shown in [Figure 4.6.2](#) on page 104) seems a good basic unit to use in the analysis. Once w_{13} is set, all of the cable weights follow. Second, the weight on a given cable is a function not only of the angle between it and its strut but also the angle between its neighbor cable and their shared strut. That will turn out to be significant later.

Now both [Equation \(4.7\)](#) and [Equation \(4.9\)](#) on the preceding page claim to be able to tell us what w_{24} is. Suppose we call the [Equation \(4.7\)](#) answer w_{24} and the [Equation \(4.9\)](#) answer $\hat{w}_{24}$. How do they relate?

$$\begin{aligned} \hat{w}_{24} \vec{e}_{24} &= w_{14} \vec{e}_{14} + w_{34} \vec{e}_{34} && \text{(Equation (4.9))} \\ &= w_{13} \vec{e}_{13} - w_{12} \vec{e}_{12} + w_{23} \vec{e}_{23} - w_{13} \vec{e}_{13} && \text{(Equation (4.6) \& Equation (4.8))} \\ &= w_{23} \vec{e}_{23} - w_{12} \vec{e}_{12} = w_{24} \vec{e}_{24} && \text{(Equation (4.7))} \end{aligned}$$

That’s rather exciting news.¹ That compatibility shows that we can establish a strictly positive stress on any such convex quadrilateral.

Suppose that we have any convex curve of vertices. If we put struts between every pair of antipodal (by arc length) points and cables between every pair of points whose distance is $1/4$ of the length of the curve, the resulting tensegrity is covered by convex quadrilaterals and so it has a strictly positive stress and is bar equivalent.

¹News that can be established also using the law of sines on the quadrilateral.

But there are two hypotheses that seem like they could be removed. First, surely $1/4$ isn't the only fraction of the curve length that works. Secondly, is it really necessary to have a convex curve?

4.6 OTHER FRACTIONS

Let's start by trying other fractions. We started with $1/4$ because the quadrilateral is pretty straightforward. Of course, if we use the fraction $1/2n$ for any $n > 1$, we get a closed $2n$ -gon of cables with struts connecting opposite points. As before, we certainly need each strut to fall in the cone generated by its cables. Also as before, we can generate families of positive solutions for each “double arrow”. We need only establish that those solutions can be compatible with each other. [Figure 4.6.1](#) shows a curve with one of the hexagonal tensegrities coming from a skip of $1/6$ of the curve length.

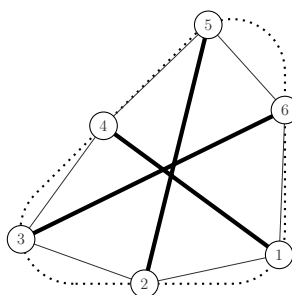


Figure 4.6.1: A more general hexagon.

Now suppose that we have two “adjacent” struts (shown in [Figure 4.6.2](#) on the following page) along with their associated cables. These two struts will share one cable on each end. Certainly we can scale the answers for the two “double arrows” so that the weights match on one of the shared cables, but it's not at all clear that we can make both ends match. And it turns out that in general it doesn't work. [Figure 4.6.3](#) on the next page shows a set of vectors that increase all strut lengths and decrease all cable lengths for this particular example, and since none of the angles are straight,

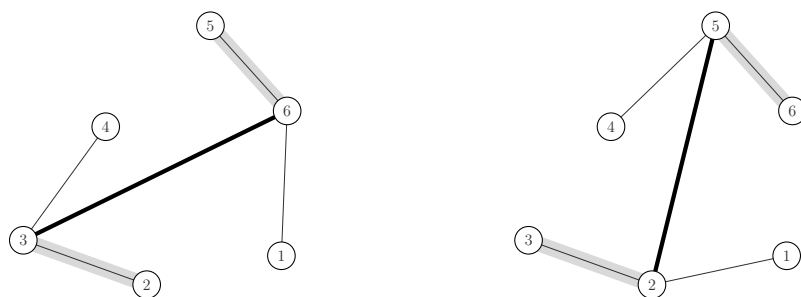


Figure 4.6.2: Two of the struts from the hexagonal tensegrity above, along with their associated cables. The two cables they share are shaded.

there will be an open set of neighbor hexagons that also have a motion. So this tensegrity is not bar equivalent.

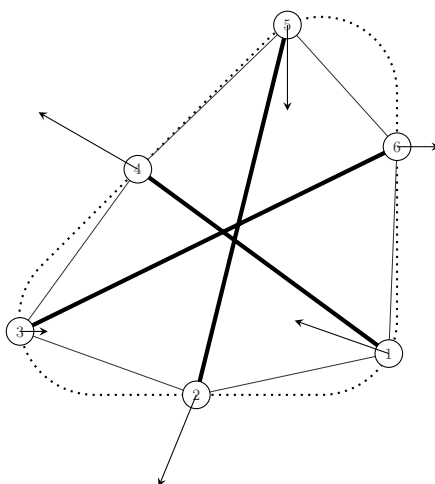


Figure 4.6.3: The hexagon with a motion.

What does work? Well, certainly any *regular* figure (where the construction at each vertex is identical to that at all of the others). [Figure 4.6.4](#) on the following page shows two regular figures that could arise.

But other figures can provide bar equivalence. Suppose, for example, that we have the non-regular hexagon shown in [Figure 4.6.5](#) on the next page. It turns out that a stress of 2 units on each of the cables and 1 unit each on the struts will put this in equilibrium, so it *is* bar equivalent, and this suggests a family of bar equivalent

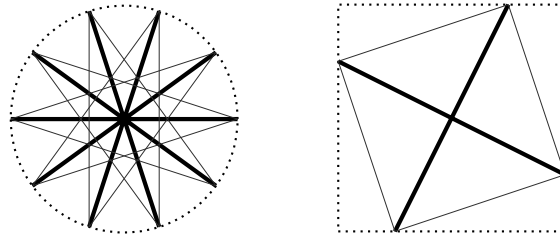


Figure 4.6.4: Two regular figures that arise with this construction.

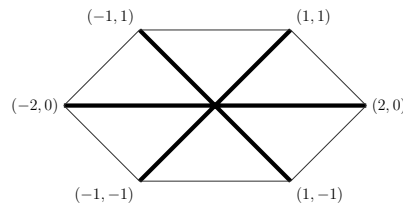


Figure 4.6.5: A bar equivalent tensegrity that is a nonregular hexagon.

hexagons. [Figure 4.6.6](#) on the following page gives another, with the weights shown on the edges: Now these are simply example hexagons. We have not addressed how much stranger they can get, whether they can actually arise while building tensegrities from curves in that fashion, nor whether, if they should arise, there must also be non-bar-equivalent subtensegrities in the same construction. This seems a field worth more exploration in the future. For now, let's see if we can settle something about non-convexity.

4.7 NON-CONVEXITY

[Figure 4.7.1](#) on page 107 shows two non-convex vertex curves with antipodal struts and fixed-skip cables. The skip in the first case is $1/4$ of the curve length, and in the second it is $1/6$. The resulting tensegrity is, in each case, covered with regular sub-

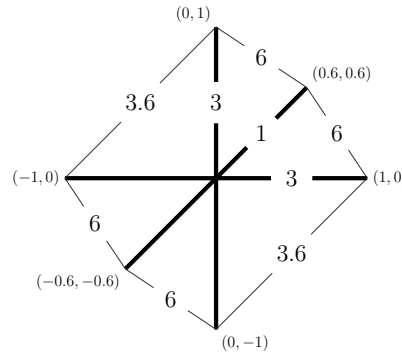


Figure 4.6.6: Another in the family of nonregular, bar equivalent hexagons.

tensegrities, all of which are bar equivalent. So each tensegrity is itself bar equivalent (though neither is infinitesimally rigid, since a vector field with the same rotational symmetry will move the subtensegrities with respect to each other).

What’s going on here? As we have seen above, what we need is for the two cables meeting at each vertex of each subtensegrity to form a cone that contains the strut. That will happen if, for a given vertex, the strut intersects the line that connects the two “distant” cable ends (the dashed lines in the [Figure 4.7.1](#) and [Figure 4.7.2](#)). Or to put it differently, if the curve has length L and the cables connect points αL apart, then the tensegrity will be bar equivalent if all segments connecting points $2\alpha L$ apart lie entirely within the curve (see [Figure 4.7.2](#) on the following page).

For a given non-convex curve, then, there is a minimum distance below which the cables must not go. This distance can be found in the following fashion. For each vertex v , have two points start at the antipode and move, one in each direction at equal speeds around the curve toward v .

If, at some point, the segment between them no longer lies entirely within the curve, stop. Call that distance d_v . Otherwise, $d_v = 0$. Now, to get a bar equivalent tensegrity, the cable skip distance must be strictly greater than $\max_{v \in \mathcal{V}} d_v$.

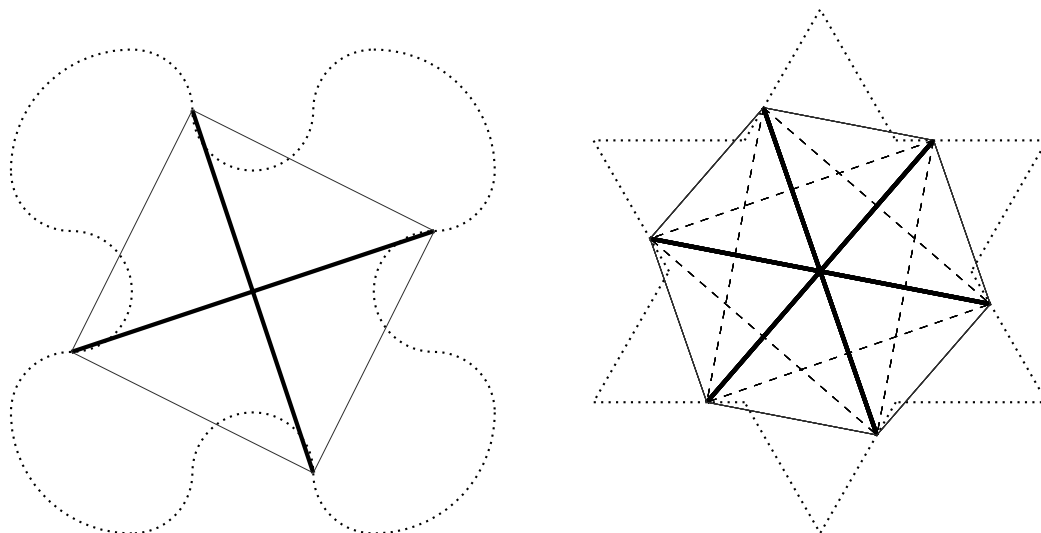


Figure 4.7.1: Here are two examples of bar-equivalent tensegrities that arise from nonconvex curves. One subtensegrity is shown for each. The dashed lines are explained in the text.

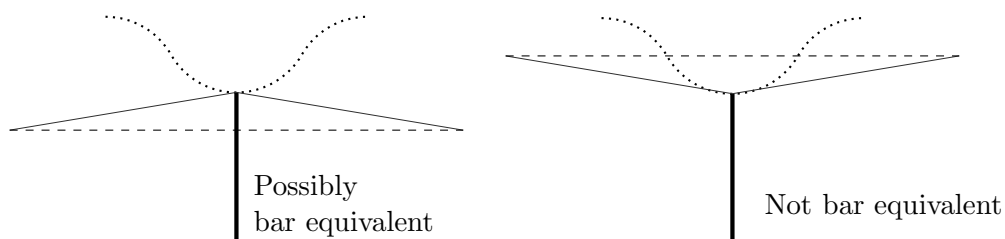


Figure 4.7.2: How the relationship between the vertex curve and the “double skip” tells on the bar equivalence of the subtensegrities.

This is not always possible. [Figure 4.7.3](#) on the next page shows a vertex curve with one of its struts. If cables are attached with a skip *anything* less than $1/2$ of the overall length, the two associated cables will both end on the same side of the strut, so the strut will not lie in the cone they create.

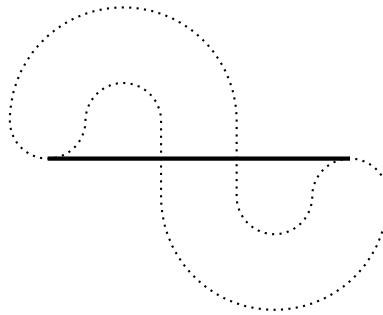


Figure 4.7.3: A vertex curve and a strut that cannot be balanced by fixed-skip cables.

4.8 AFFINE TRANSFORMATIONS

4.8.1 DEEPER INTO THE CROSSED SQUARE

In [Section 2.2](#) on page [24](#), we found that the crossed square (also shown in [Figure 4.8.1\(a\)](#)) is bar equivalent. Let's find how its set of loads, $Y(\mathcal{X})$, looks.

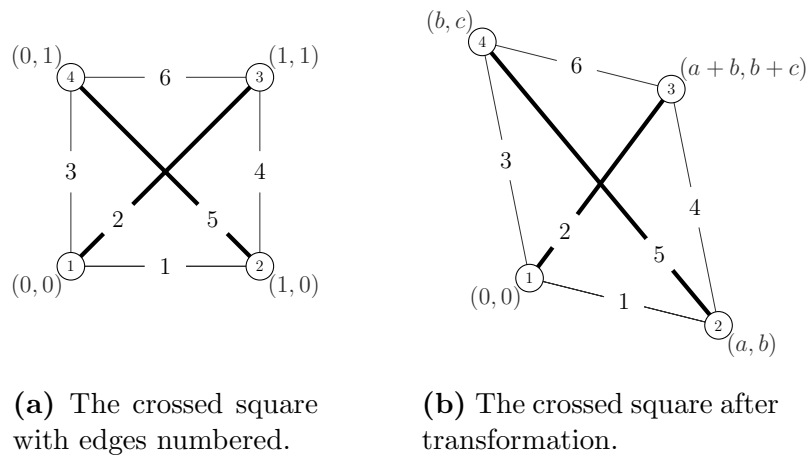


Figure 4.8.1: The crossed-square example, both the original and after being transformed by the matrix $L^\top L = \begin{bmatrix} a & b \\ b & c \end{bmatrix}$.

Suppose $V = (x_1, y_1, x_2, y_2, x_3, y_3, x_4, y_4) \in \text{VF}(\mathcal{V})$. Then, letting the edges be e_1 through e_6 (as numbered in the figure), we have

$$\begin{aligned} YV(e_1) &= -(x_2 - x_1) & YV(e_4) &= -(y_3 - y_2) \\ YV(e_2) &= x_3 - x_1 + y_3 - y_1 & YV(e_5) &= x_2 - x_4 - (y_2 - y_4) \\ YV(e_3) &= -(y_4 - y_1) & YV(e_6) &= -(x_3 - x_4) \end{aligned}$$

But then

$$\begin{aligned} \sum_{i=1}^6 YV(e_i) &= (-x_2 + x_1) + (x_3 - x_1 + y_3 - y_1) + (-y_4 + y_1) \\ &\quad + (-y_3 + y_2) + (x_2 - x_4 - y_2 + y_4) + (-x_3 + x_4) = 0 \end{aligned}$$

So $\text{im } Y$ is contained in the hyperplane $\sum_{i=1}^6 YV(e_i) = 0$. On the other hand, given any six real numbers $r_1, \dots, r_6$ with $r_1 + \dots + r_6 = 0$, the vector field

$$(r_1, r_3, 0, r_4, r_1 + r_2 + r_3, 0, -r_4 - r_5, 0)$$

is one of (infinitely) many for which $YV(e_i) = r_i, i \in \{1, \dots, 6\}$.

So $\text{im } Y$ is the hyperplane $\sum_{i=1}^6 YV(e_i) = 0$ and from that equation, it is clear to see that $\text{im } Y \cap C(\mathcal{E})^+ = \{\mathbf{0}\}$ (whenever one of the $YV(e_i)$ is positive, another must be negative to offset it). We note that the hyperplane is 5-dimensional, a reasonable thing since $\text{VF}(\mathcal{V})$ is 8-dimensional and there are three degrees of freedom given to rigid motions of the plane (one rotational and two translational) that, of course, lie in the kernel of Y .

Now suppose that we apply the linear transformation whose matrix is L to our example. How will $\text{im } Y$ change? Well, Y is defined by taking the dot product of

vectors. We note that

$$\begin{aligned}
 (LV(v_1) - LV(v_2)) \cdot (Lp(v_1) - Lp(v_2)) \\
 &= L(V(v_1) - V(v_2)) \cdot L(p(v_1) - p(v_2)) \\
 &= (V(v_1) - V(v_2))^T L^T L(p(v_1) - p(v_2)) \\
 &= (V(v_1) - V(v_2)) \cdot L^T L(p(v_1) - p(v_2)).
 \end{aligned} \tag{4.12}$$

So we can accomplish the same purpose by transforming p by $L^T L$ and leaving V alone. A simple calculation shows that $L^T L$ is always symmetric (with nonnegative entries on the diagonal), so we'll write it as

$$L^T L = \begin{bmatrix} a & b \\ b & c \end{bmatrix}$$

and see what effect it has (one possibility is shown in [Figure 4.8.1\(b\)](#)).

Now we get (using $\hat{Y}$ for the rigidity operator after transformation)

$$\begin{aligned}
 \hat{Y}V(e_1) &= -a(x_2 - x_1) - b(y_2 - y_1) \\
 \hat{Y}V(e_2) &= (a + b)(x_3 - x_1) + (b + c)(y_3 - y_1) \\
 \hat{Y}V(e_3) &= -b(x_4 - x_1) - c(y_4 - y_1) \\
 \hat{Y}V(e_4) &= -b(x_3 - x_2) - c(y_3 - y_2) \\
 \hat{Y}V(e_5) &= (a - b)(x_2 - x_4) + (b - c)(y_2 - y_4) \\
 \hat{Y}V(e_6) &= -a(x_3 - x_4) - b(y_3 - y_4).
 \end{aligned}$$

The arithmetic is only touch more complicated, but the result is the same, $\text{im } Y$ is the hyperplane $\sum_{i=1}^6 YV(e_i) = 0$. So the crossed square remains bar equivalent under linear transformation.

This is far from obvious. Linear transformation does *not*, in general, preserve dot product. So a vector field that is in the kernel of Y before transformation is not necessarily in the kernel of $\hat{Y}$ afterward (see [Figure 4.8.2](#) on the next page for an example). This bears more exploration.

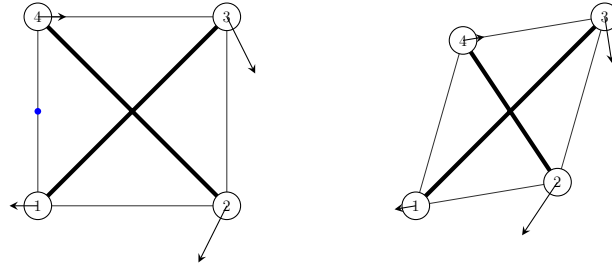


Figure 4.8.2: The crossed-square, original and transformed, with a motion (rotation around the marked point) that is in $\ker Y$ before transformation by $L = \begin{bmatrix} 3/4 & 1/4 \\ 1/8 & 7/8 \end{bmatrix}$ and not in $\ker \hat{Y}$ afterward.

4.8.2 TRIANGLES AGAIN

Let's look at one more example and see if we can figure out what is going on. In Figure 4.8.3 we find a tensegrity that looks like the one in Figure 3.5.2 on page 68, but we are going to treat it differently. This time we are going to define $\mathcal{X}$ to be those

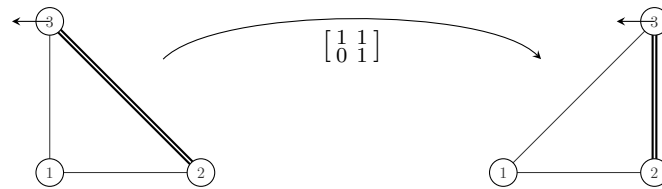


Figure 4.8.3: A tensegrity that is $\mathcal{X}$ -bar equivalent before transformation by $\begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$ but not afterward. The motion shown, moving only vertex 3, results in a load that has negative components for the lefthand tensegrity, but is semipositive for the right.

variations that do not change the length of the cables (to first order). That means

$$\mathcal{X} = \{(x_1, y_1, x_2, y_2, x_3, y_3) \in \text{VF}(\mathcal{V}) : x_1 = x_2, y_1 = y_3\}$$

The bar functions as a strut and a cable, so $C(\mathcal{E})$ is a 4-dimensional space. Taking the edges in the order “1-2 cable, 1-3 cable, 2-3 cable, 2-3 strut”, the elements of $Y(\mathcal{X})$ turn out to all be of the form

$$(0, 0, -(y_3 - y_2) - (x_2 - x_3), (y_3 - y_2) + (x_2 - x_3))$$

so $Y(\mathcal{X}) \cap C(\mathcal{E})^+ = \{\mathbf{0}\}$ and the tensegrity is $\mathcal{X}$ -bar equivalent.

On the other hand, if we transform the tensegrity with $L = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$, then we get

$$L\mathcal{X} = \{(x_1, y_1, x_2, y_2, x_3, y_3) \in \text{VF}(\mathcal{V}) : x_1 + y_1 = x_2 + y_2, y_1 = y_3\}$$

and so the elements of $\hat{Y}(L\mathcal{X})$ are

$$(y_2 - y_1, x_1 - x_3, y_1 - y_2, y_2 - y_1).$$

Now $\hat{Y}(L\mathcal{X}) \cap C(\mathcal{E})^+$ contains functions that are zero except on the “newly diagonal” cable, the edge running from vertex 1 to vertex 3. On that edge, they can be strictly positive.

Since there is a semipositive stress, namely $(1, 0, 2, 1)$, our transformed tensegrity is partially $L\mathcal{X}$ -bar equivalent, but it is not fully so.

4.8.3 CONCLUSION

So what is the difference between the two preceeding examples? $\mathcal{X}$.

To be a little clearer, let's think back to [Equation \(4.12\)](#) on page 110 in which we showed that for two vectors v and p , $Lv \cdot Lp = v \cdot L^\top Lp$. We could have gone the other direction and we would have gotten

$$(LV(v_1) - LV(v_2)) \cdot (Lp(v_1) - Lp(v_2)) = L^\top L(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2)).$$

That equation gives the answer to the question. In our first example, $L^\top L$ was an automorphism of $\mathcal{X}$ (which, after all, was all of $\text{VF}(\mathcal{V})$). In the second it was not. So that suggests a proposition.

Proposition 4.8.1. *Let $G(p)$ be an $\mathcal{X}$ -bar equivalent tensegrity in $\mathbb{R}^n$, L a linear transformation on $\mathbb{R}^n$ and $t \in \mathbb{R}^n$. Define A by $Ax = Lx + t$ for all $x \in \mathbb{R}^n$. Then $G(Ap)$ is $A\mathcal{X}$ -bar equivalent if $L^\top L\mathcal{X} \subset \mathcal{X}$.*

Proof. Let $AV \in A\mathcal{X}$. Then, for any $\{v_1, v_2\} \in \mathcal{E}$ we have

$$\begin{aligned}
\hat{Y}(AV)(\{v_1, v_2\}) &= \pm(LV(v_1) + t - LV(v_2) - t) \cdot (Lp(v_1) + t - Lp(v_2) - t) \\
&= \pm(LV(v_1) - LV(v_2)) \cdot L(p(v_1) - p(v_2)) \\
&= \pm L^\top (LV(v_1) - LV(v_2)) \cdot (p(v_1) - p(v_2)) \\
&= \pm(L^\top LV(v_1) - L^\top LV(v_2)) \cdot (p(v_1) - p(v_2)) \\
&= Y(L^\top LV)(\{v_1, v_2\})
\end{aligned}$$

(where the choice of sign depends on the type of edge). But since $L^\top L\mathcal{X} \subset \mathcal{X}$, we have $\hat{Y}(AV) \in Y(\mathcal{X})$. So $\hat{Y}(A\mathcal{X}) \subset Y(\mathcal{X})$, which means that

$$\hat{Y}(A\mathcal{X}) \cap C(\mathcal{E})^+ \subset Y(\mathcal{X}) \cap C(\mathcal{E})^+ = \{\mathbf{0}\}$$

and $G(Ap)$ is $A\mathcal{X}$ -bar equivalent.

Note that the proposition and proof still stand if “ $\mathcal{X}$ -bar equivalent” is replaced by “partially $\mathcal{X}$ -bar equivalent”, “ $C(\mathcal{E})^+$ ” is replaced with “ $\text{int } C(\mathcal{E})^+$ ” and “ $\{\mathbf{0}\}$ ” by “ $\emptyset$ ”. $\square$

Remembering that “bar equivalent” means “ $\text{VF}(\mathcal{V})$ -bar equivalent”, we immediately get a corollary.

Corollary 4.8.2. *If $G(p)$ is a bar equivalent tensegrity in $\mathbb{R}^n$ and A is an affine transformation on $\mathbb{R}^n$, then $G(Ap)$ is bar equivalent.*

Proof. Note that for any linear transformation L , $L^\top LVF(\mathcal{V}) \subset VF(\mathcal{V})$, and apply Proposition 4.8.1. $\square$

In the case of an invertible transformation (for which we only need that L is invertible, since $A^{-1}x = L^{-1}(x - t)$), we can say something more.

Proposition 4.8.3. *Let $G(p)$ be an $\mathcal{X}$ -bar equivalent tensegrity in $\mathbb{R}^n$, L an invertible linear transformation and $t \in \mathbb{R}^n$. Define A by $Ax = Lx + t$. Then $G(Ap)$ is $(L^\top)^{-1}\mathcal{X}$ -bar equivalent.*

Proof. Let $(L^\top)^{-1}V \in (L^\top)^{-1}\mathcal{X}$. Then, for any $\{v_1, v_2\} \in \mathcal{E}$, we have

$$\begin{aligned}
 \hat{Y}((L^\top)^{-1}V)(\{v_1, v_2\}) &= ((L^\top)^{-1}V(v_1) - (L^\top)^{-1}V(v_2)) \cdot (Ap(v_1) - Ap(v_2)) \\
 &= (L^\top)^{-1}(V(v_1) - V(v_2)) \cdot (Lp(v_1) + t - Lp(v_2) - t) \\
 &= (L^\top)^{-1}(V(v_1) - V(v_2)) \cdot L(p(v_1) - p(v_2)) \\
 &= L^\top (L^\top)^{-1}(V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2)) \\
 &= (V(v_1) - V(v_2)) \cdot (p(v_1) - p(v_2)) \\
 &= YV(\{v_1, v_2\})
 \end{aligned}$$

So $\hat{Y}((L^\top)^{-1}V) = YV \in Y(\mathcal{X})$ and we already know that $Y(\mathcal{X}) \cap C(\mathcal{E})^+ = \{\mathbf{0}\}$.

Once again, the proposition and proof are true, *mutatis mutandis* for partially $(L^\top)^{-1}\mathcal{X}$ -bar equivalent. $\square$

4.8.4 MORE CIRCLES

So what does this say about the on-a-circle example of [Section 4.2](#)? Since $\mathcal{X} = \text{VF}(\mathcal{V})$, [Corollary 4.8.2](#) says that any affine transformation of the on-a-circle example is still bar equivalent. Note that this is different from what one would get by first transforming the circle of vertices and then applying the on-a-circle construction (see [figure Figure 4.8.4](#)).

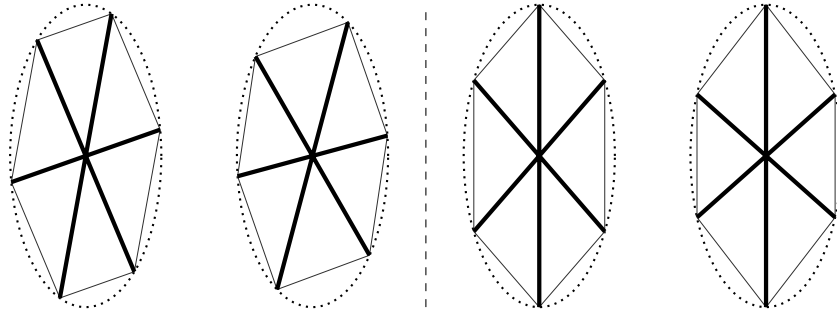


Figure 4.8.4: Two ways to build elliptical tensegrities. Two different subtensegrities have been shown. In both cases the picture on the left is the result of building a circular tensegrity and then transforming it and the picture on the right is the result of building the tensegrity on an ellipse.

4.9 A CORNER CASE

What might it look like to have an example where not every subtensegrity is bar equivalent? Let's try to build such a creature. Consider the family of tensegrities represented in [Figure 4.9.1](#). For the first triangle (where each side is length 2), a stress

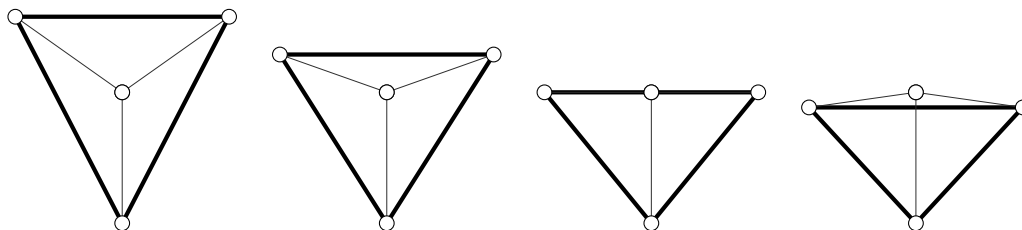


Figure 4.9.1: Representatives of our family of triangular tensegrities.

which has weight 3 on all of the cables and weight 1 on each of the struts will suffice to show that it is bar equivalent. As the top comes down to meet the midpoint, the weights on the upper two cables must rise (or all of the others fall) until, when the top strut either touches or runs below the midpoint, the tensegrity is no longer bar equivalent.

Let's group together a lot of the bar equivalent ones and slip one that isn't bar equivalent in amongst them. We can choose a family of such tensegrities with the top strut approaching the center point, as shown in [Figure 4.9.2](#) (we've removed the struts from the picture for visibility's sake).

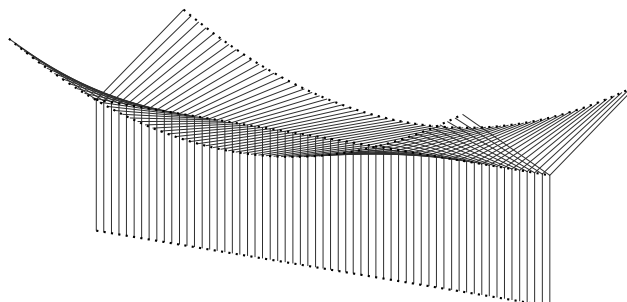


Figure 4.9.2: Triangle subtensegrities with varying top angle forming a continuous tensegrity. The outer strut triangles are not shown.

Now in three dimensions, even those subtensegrities that previously were infinitesimally rigid are no longer so, as they have no way to resist variations normal to their planes of definition. However, as Proposition 4.9.1 tells us, they remain bar equivalent.

Proposition 4.9.1. *A finite tensegrity that is bar equivalent in $\mathbb{R}^n$ is bar equivalent in Euclidean space of all dimensions higher than n .*

Proof. Let $G(p)$ be a finite, bar equivalent tensegrity in $\mathbb{R}^n$. By Lemma 2.4.3 on page 35, that tensegrity has a strictly positive stress, that is, a positive linear dependence among the row vectors of the rigidity matrix. Moving to a higher dimension will change those vectors, but only by adding columns of zeros for the new dimension, so it doesn't change that linear dependence. Since $G(p)$ continues to have a positive stress, it continues to be bar equivalent. $\square$

So our continuous family is countably covered by bar equivalent subtensegrities and hence must be bar equivalent itself. That means, of course, that $Y(\mathcal{X}) \cap C(\mathcal{E})^+ = \{\mathbf{0}\}$, which is to say, no matter which V we pick, YV must have some negative values.

That's true, but there's something a little strange going on here. It turns out that the image of Y , while intersecting the nonnegative orthant only at the origin, gets arbitrarily close to it elsewhere. Let's see if we can watch that happen.

Take the vertical cable in each subtensegrity. We can shrink this cable by putting a length-1 vector pointing directly downward at the central node and balancing it with a length $1/3$ upward-pointing vector at each of the other three nodes (see Figure 4.9.3 on the following page).

This set of vectors will have no effect whatsoever on the struts. It will produce a value of $4/3$ on the vertical cable, and it will produce a force on the other two cables that depends on their angle (see Figure 4.9.4 on the next page).

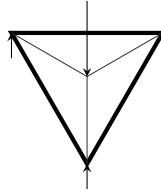


Figure 4.9.3: Balancing forces placed on a subtensegrity.

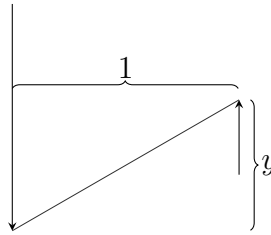


Figure 4.9.4: The forces present on one “arm” of the subtensegrity.

In every case, the horizontal length of each of those cables is 1. If the vertical length is y , then the rigidity operator will take our set of vectors to the value $-2/3 \frac{y}{\sqrt{1+y^2}}$. Now when $y = 0$, this value is 0. That makes sense, since the forces are now acting orthogonally to the edge, but we can use that fact to our advantage.

Suppose, now, that we identify our family of subtensegrities with the interval $[-1, 1]$ in such a way that if $x \in [-1, 1]$ gives a certain subtensegrity, then $y = x^2$ for that tensegrity (this is function used in [Figure 4.9.2](#)). Now we take a family of vector fields on our vertices. For every value $i \in \mathbb{Z}$, we let $V_i(x)$ be the zero vector for $x \notin [-\frac{1}{i}, \frac{1}{i}]$ and then have vectors in the ratios we’ve discussed increasing linearly from 0 at $x = -\frac{1}{i}$ to 1 at $x = 0$ and back down to 0 at $x = \frac{1}{i}$ (the case $i = 2$ is shown in [Figure 4.9.5](#) on the following page).

Now YV_i is a function that is strictly positive on the center subtensegrity. After all, it shrinks that center cable, and since it acts orthogonally to them, it doesn’t

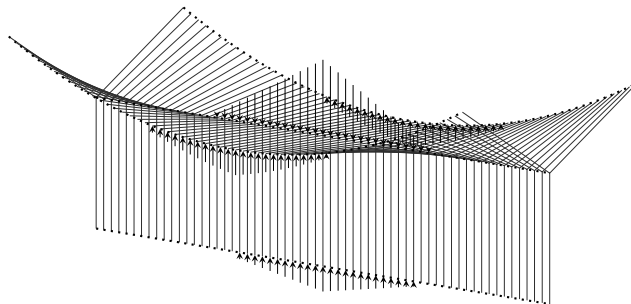


Figure 4.9.5: The tensegrity with the vector field described.

change the lengths of the other cables, to first order. Of course, YV_i is negative on the upper cables of those tensegrities surrounding the center one—we knew it had to be negative somewhere since the tensegrity is bar equivalent.

For each V_i , we can define $f_i(e) = \max\{YV_i(e), 0\}$. Then $f_i \in C(\mathcal{E})^+$, and as i gets large, the distance $\|YV_i - f_i\|$ gets arbitrarily small. On the other hand $\|YV_i\| = 1 = \|f_i\|$ for all i , so we have a subspace, $\text{im } Y$ that, at a distance 1 from the origin, gets arbitrarily close to the nonnegative orthant and yet which intersects that orthant only at the origin.

CHAPTER 5

CONCLUSION AND FUTURE DIRECTIONS

So where are we and where do we go from here?

We set out to take the theorem “A tensegrity is bar equivalent if and only if it has a strictly positive stress” into the new world of continuous tensegrities. We’ve come close to accomplishing that. We have:

- i) A tensegrity is partially bar equivalent if and only if it has a semipositive stress.
- ii) A tensegrity is bar equivalent if it has a strictly positive stress.
- iii) A tensegrity is minimally bar equivalent only if it has a strictly positive stress.
- iv) If a tensegrity is countably covered with subtensegrities that have strictly positive stresses, then it has a strictly positive stress.

If we want to show that a tensegrity has a motion, showing that it has no semipositive stress gives a very strong result. On the other hand, if we wish to show that a tensegrity has no motion, showing that it has a strictly positive stress gives a strong result and showing that it has a semipositive stress still gives a valuable result. We can handle all of these situations — but only when $\mathcal{E}$ is compact.

That is not always a given. For example, Connelly et al. (2003) find a flow of any nonconvex polygon in $\mathbb{R}^2$ that “convexifies” the polygon and in which all non-adjacent vertices are constantly moving away from each other. Moving into the continuous realm (see [Figure 5.0.1](#) on the next page), we find that for each vertex, we get a “fan” of struts.

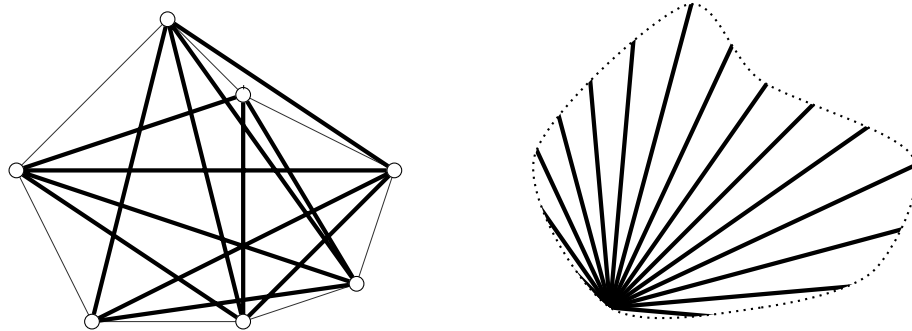


Figure 5.0.1: Moving the “convexifying polygons” problem in to the continuous world. In the continuous picture, the struts from only one vertex are shown.

Within that fan, we can find a sequence of struts which limits to the “zero-length strut” connecting the vertex to itself. If we exclude those “zero-length struts”, the edgeset is not compact. If we include them, the load on those struts will have to be zero, so any motion is semipositive at best. On the other hand, perhaps it could be shown that the motion, though semipositive, is strictly positive on all of the struts connecting distinct points.

In the work here, we have not addressed the question of infinitesimal rigidity. A bar equivalent continuous tensegrity is infinitesimally rigid if the associated bar framework is, but determining whether a continuous bar framework is infinitesimally rigid is not always simple. Roth (1981) gives a “rigidity predictor” for bar frameworks. Perhaps it could be extended to continuous bar frameworks and thence to continuous tensegrities.

Any work on infinitesimal rigidity of continuous bar frameworks is likely to be informed by the literature on continuous families of linear equalities. That is a realization which may well be exceedingly fruitful. What we have done here is valuable in its own right in solving problems with continuous tensegrities, but the statement that $YV \in C(\mathcal{E})+$ is really just an infinite family of linear inequalities. By modi-

fying Y , the theory could be applied to other situations. For example, Ashton et al. (in revision) use both struts and kinks (places where the vertex curve has reached a curvature constraint) to build the rigidity operator. A similar idea could perhaps be employed to provide an alternate proof of Schur’s Theorem (page 6).

As another example, consider moving the differential version of the “convexifying” task up one dimension. In this realm, we could still seek to have all vertex-vertex distances increase, or we could instead build our rigidity operator using the volumes of the tetrahedra determined by sets of 4 points on the surface, and we could require that those volumes be strictly increasing. Or, with careful choice of rigidity operator, perhaps our results can provide new proofs for various of the Theorems of the Alternative (see Section 3.1 and Appendix B).

Speaking of Theorems of the Alternative, at the end of Appendix B, we relate stresses of a tensegrity to variations lying in $\mathcal{X}^\perp$. It would be interesting to know how those variations relate to the “curvature force” of Cantarella et al. (2006). For that matter, there is plenty of work yet to be done to understand exactly which variations thus arise. In fact, at this point, the circle-of-struts example is the only bar equivalent tensegrity we have whose design variations are local isometries of the vertex curve.

In the world where all variations are design variations, there is exploration that could be done into which tensegrities that arise from the antipodal strut/fixed-skip cable technique are bar equivalent (see Section 4.6 on page 105). Also, the continuous analog of the Grünbaum polygons from page 6 waits to be found.

Finally, there are those conjectures from pages 87 and 91:

Conjecture 3.8.4. *A tensegrity is minimally $\mathcal{X}$ -bar equivalent if and only if it admits only one semipositive stress, up to scaling. Furthermore, that stress is strictly positive.*

Conjecture 3.9.3. *Every $\mathcal{X}$ -bar equivalent continuous tensegrity is countably covered by minimally $\mathcal{X}$ -bar equivalent subtensegrities.*

Conjecture 3.9.4. *Every $\mathcal{X}$ -bar equivalent tensegrity has a strictly positive stress.*

With regard to that last one, having the strictly positive loads in the interior of $C(\mathcal{E})^+$ and the semipositive ones on the boundary seems appropriate but that $C^*(\mathcal{E})^+$ has no interior is surprising. Perhaps a different topology on $C(\mathcal{E})$ would give an interior to $C^*(\mathcal{E})^+$ and place the strictly positive stresses in that interior and the semipositive ones along the boundary. Or perhaps one of the other Theorems of the Alternative, such as those of Tucker (1956) or Slater (1951) could solve the problem with the topology as it stands.

BIBLIOGRAPHY

- Artin, Michael. 1991. *Algebra*, Prentice Hall Inc., Englewood Cliffs, NJ. MR **1129886** (**92g**:00001)
- Ashton, Ted, Jason Cantarella, Michael Piatek, and Eric Rawdon. *Knot Tightening by Constrained Gradient Descent*, in revision.
- Burago, Dmitri, Yuri Burago, and Sergei Ivanov. 2001. *A Course in Metric Geometry*, Graduate Studies in Mathematics, vol. 33, American Mathematical Society, Providence, RI. MR **1835418** (**2002e**:53053)
- Cantarella, Jason, Joseph H.G. Fu, Robert B. Kusner, John M. Sullivan, and Nancy Wrinkle. 2006. *Criticality for the Gehring Link Problem*, Geometry & Topology **10**, 2055–2115, available at [arXiv:math.DG/0402212](https://arxiv.org/abs/math/0402212).
- Chen, Christopher S. and Donald E. Ingber. 1999. *Tensegrity and Mechanoregulation: from Skeleton to Cytoskeleton*, Osteoarthritis and Cartilage **7**, 81–94, available at <http://www.idealibrary.com>.
- Connelly, Robert. 1982. *Rigidity and Energy*, Invent. Math. **66**, no. 1, 11–33. MR **652643** (**83m**:52012)
- Connelly, Robert, Erik D. Demaine, and Günter Rote. 2003. *Straightening Polygonal Arcs and Convexifying Polygonal Cycles*, Discrete & Computational Geometry **30**, 205–239.
- Conway, John B. 1990. *A Course in Functional Analysis*, Second edition, Graduate Texts in Mathematics, vol. 96, Springer-Verlag, New York. MR **1070713** (**91e**:46001)
- Covitz, H. and S. B. Nadler Jr. 1970. *Multi-valued Contraction Mappings in Generalized Metric Spaces*, Israel J. Math. **8**, 5–11. MR 0263062 (41 #7667)

- Craven, Bruce Desmond. 1978. *Mathematical Programming and Control Theory*, Chapman and Hall, London. Chapman and Hall Mathematics Series. MR **515723** (80d:93001)
- do Carmo, Manfredo P. 1976. *Differential Geometry of Curves and Surfaces*, Prentice-Hall, Upper Saddle River, NJ.
- Eaves, Elisabeth. 2004. *Art for Smart People*, Slate, available at <http://www.slate.com/id/2093711/>.
- Folland, Gerald B. 1999. *Real Analysis*, Second edition, Pure and Applied Mathematics (New York), John Wiley & Sons Inc., New York. Modern techniques and their applications; A Wiley-Interscience Publication. MR **1681462** (2000c:00001)
- Forest Products Laboratory. 1999. *Wood Handbook: Wood as an Engineering Material*, United States Department of Agriculture, Madison, WI, General Technical Report FPL-GTR-113.
- Farrell, H. M., Jr., E. M. Brown, P. D. Hoagland, and E. L. Malin. *Higher Order Structures of the Caseins: a Paradox?*, Advanced Dairy Chemistry. A: Proteins, Third edition (P. F. Fox and Paul McSweeney, eds.), Vol. 1, Kluwer Academic/Plenum Publishers, New York, 2003.
- Franklin, Eric. 1996. *Dynamic Alignment through Imagery*, Human Kinetics, Champaign, IL.
- Gaal, Steven A. 1973. *Linear Analysis and Representation Theory*, Springer-Verlag, New York. Die Grundlehren der mathematischen Wissenschaften, Band 198. MR 0447465 (56 #5777)
- Girvin, Harvey F. 1944. *Strength of Materials*, First edition (Charles Edward O'Rourke, ed.), International Texts in Civil Engineering, International Textbook Company, Scranton, Pennsylvania.

- Gordan, Paul. 1873. *Ueber die Auflösung linearer Gleichungen mit reellen Coefficienten*, Mathematische Annalen **6**, 23–28 (German).
- Grun, Bernard. 1979. *The Timetables of History: A Horizontal Linkage of People and Events*, Simon and Schuster, New York.
- Grünbaum, Branko and Geoffrey C. Shephard. *Lectures in Lost Mathematics*. Unpublished lecture notes, University of Washington.
- Jeong, Byeongha, Jin-Sung Park, Kyoung J. Lee, Seok-Cheol Hong, Ju-Yong Hyon, Hyun Choi, Dong June Ahn, and Seokmann Hong. Jan. 2007. *Direct Measurement of the Force Generated by a Single Macrophage*. 1, J. Korean Phys. Soc. **50**, no. 1, 313–319.
- Kolář, Ivan, Peter W. Michor, and Jan Slovák. 1993. *Natural Operations in Differential Geometry*, Springer-Verlag, Berlin. [Electronic version, corrected; accessed 15-March-2007]. MR **1202431** (**94a**:58004)
- Landau, Edmund. 1958. *Elementary Number Theory*, Chelsea Publishing Co., New York, N.Y. Translated by J. E. Goodman. MR 0092794 (19,1159d)
- Lawson, Terry. 2003. *Topology: a Geometric Approach*, Oxford Graduate Texts in Mathematics, vol. 9, Oxford University Press, Oxford. MR **1984838** (**2004c**:57001)
- Lenarčič, Jadran and Carlo Galletti (eds.) 2004. *On Advances in Robot Kinematics*, Kluwer Academic Publishers, The Netherlands.
- Luxemburg, W. A. J. 1958. *On the Convergence of Successive Approximations in the Theory of Ordinary Differential Equations*. II, Proceedings of the Koninklijke Nederlandse Akademie van Wetenschappen. Series A, Mathematical Sciences. (now Indagationes Mathematicae). **20**, 540–546. MR 0124554 (23 #A1866)
- Maina, J. N. 2007. *Spectacularly Robust! Tensegrity Principle Explains the Mechanical Strength of the Avian Lung*, Respiratory Physiology & Neurobiology **155**, no. 1, 1-10.
- Mangasarian, Olvi L. 1969. *Nonlinear Programming*, McGraw-Hill Book Co., New York. MR 0252038 (40 #5263)

- Martin, R. Bruce, David B. Burr, and Neil A. Sharkey. 1998. *Skeletal Tissue Mechanics*, Springer-Verlag New York, Inc.
- MatWeb. 2007. *MatWeb: Material Property Data*, Automation Creations, Inc. [Online; accessed 27-March-2007].
- Morrison, Terry J. 2001. *Functional Analysis: An Introduction to Banach Space Theory*, John Wiley & Sons, inc., New York.
- Motzkin, Theodore S. 1936. *Beiträge zur Theorie der Linearen Ungleichungen*, Basel, Jerusalem, Inaugural Dissertation.
- Munkres, James R. 2000. *Topology*, Second edition, Prentice Hall, Upper Saddle River, NJ.
- Murray, Richard M., Zexiang Li, and S. Shankar Sastry. 1994. *A Mathematical Introduction to Robotic Manipulation*, CRC Press, Boca Raton, FL. MR **1300410** (95k:70010)
- Pardon, John. *On the Unfolding of Simple Closed Curves*, private communication. Work dated November 12, 2006.
- Ponnusamy, S. 2002. *Foundations of Functional Analysis*, Alpha Science International Ltd., Pangbourne. MR **1998947** (2004g:46001)
- Pugh, Anthony. 1976. *An Introduction to Tensegrity*, University of California Press, Berkeley and Los Angeles, California.
- Rockafellar, R. Tyrrell. 1970. *Convex Analysis*, Princeton University Press, Princeton, New Jersey.
- Rote, Günter, Francisco Santos, and Ileana Streinu. 2003. *Expansive Motions and the Polytope of Pointed Pseudo-Triangulations*, Discrete and Computational Geometry, Algorithms Combin., vol. 25, Springer, Berlin, pp. 699–736. MR **2038499** (2005j:52019)

- Roth, B. 1981. *Rigid and Flexible Frameworks*, Amer. Math. Monthly **88**, no. 1, 6–21.
MR **619413** (**83a**:57027)
- Roth, B. and W. Whiteley. 1981. *Tensegrity Frameworks*, Trans. Amer. Math. Soc. **265**, no. 2, 419–446. MR **610958** (**82m**:51018)
- Roth, George B. Apr. 2005. *Snoring and Sleep Apnea — Structural Implications*, American Chiropractor **27**, no. 4, 22–23.
- Rotman, Joseph J. 1995. *An Introduction to the Theory of Groups*, Fourth edition, Graduate Texts in Mathematics, vol. 148, Springer-Verlag, New York. MR **1307623** (**95m**:20001)
- Royden, H. L. 1968. *Real Analysis*, Second edition, Macmillan, New York.
- Rudin, Walter. 1987. *Real and complex analysis*, 3rd ed., McGraw-Hill Book Co., New York. MR **924157** (**88k**:00002)
- Shackelford, James F. and William Alexander (eds.) 2001. *CRC Materials Science and Engineering Handbook*, Third edition, CRC Press, Boca Raton, FL.
- Sharpe, R. W. 1997. *Differential Geometry*, with a forward by S. S. Chern, Graduate Texts in Mathematics, vol. 166, Springer-Verlag, New York. Cartan’s generalization of Klein’s Erlangen program. MR **1453120** (**98m**:53033)
- Slater, Morton L. 1951. *A Note on Motzkin’s Transposition Theorem*, Econometrica **19**, 185–187. MR 0061636 (15,857e)
- Snelson, Kenneth D. 1965. *Continuous Tension, Discontinuous Compression Structures*, U.S. Patent 3,169,611.
- Stiemke, Erich. 1915. *Über positive Lösungen homogener linearer Gleichungen*, Mathematische Annalen **76**, 340–342.
- Strichartz, Robert S. 1995. *The Way of Analysis*, Jones and Bartlett Pub. Intl., Boston, MA.

Sullivan, John M. 2006. *Curves of Finite Total Curvature*, lecture notes, available at arXiv:math/0606007v1.

Tabor, Jacek and Józef Tabor. 2002. *Geometrical Aspects of Stability*, Functional Equations—Results and Advances, Adv. Math. (Dordr.), vol. 3, Kluwer Acad. Publ., Dordrecht, pp. 123–132. MR **1912709** (2003e:41045)

Tucker, A. W. 1956. *Dual Systems of Homogeneous Linear Relations*, Linear Inequalities and Related Systems, Annals of Mathematics Studies, no. 38, University Press, Princeton, N. J., pp. 3–18. MR 0089112 (19,621a)

Volokh, K. Yu., O. Vilnay, and M. Belsky. *Cell Cytoskeleton and Tensegrity*, Mechanobiology: Cartilage and Chondrocyte (J.-F. Stoltz, ed.), Vol. 2, IOS Press, 2002, pp. 63–67.

Wikipedia. 2007a. *Prestressed concrete* — *Wikipedia, The Free Encyclopedia*, available at en.wikipedia.org/w/index.php?title=Prestressed_concrete&oldid=110868200. [Online; accessed 28-February-2007].

———. 2007b. *Metric (mathematics)* — *Wikipedia, The Free Encyclopedia*, available at [en.wikipedia.org/w/index.php?title=Metric_\(mathematics\)&oldid=116734209](http://en.wikipedia.org/w/index.php?title=Metric_(mathematics)&oldid=116734209). [Online; accessed 3-April-2007].

———. 2006. *Support (measure theory)* — *Wikipedia, The Free Encyclopedia*, available at [en.wikipedia.org/w/index.php?title=Support_\(measure_theory\)&oldid=95917825](http://en.wikipedia.org/w/index.php?title=Support_(measure_theory)&oldid=95917825). [Online; accessed 29-March-2007].

APPENDIX A

EUCLIDEAN MOTIONS: INFINITESIMAL RIGID MOTIONS OF SPACE

A.1 THE SPECIAL ORTHOGONAL GROUP: SO_n AND $\mathfrak{so}_n$

As we mentioned on page 17, the vector fields in $\text{VF}(\mathcal{V})$ vary in nature. Figure A.1.1 shows three different vector fields on the same tensegrity. We are here interested in identifying which vector fields, like the first two of those in the figure, can arise from rigid motions of space.

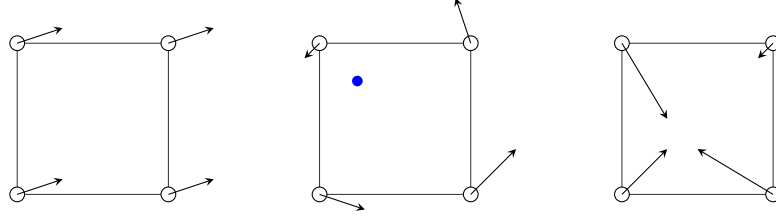


Figure A.1.1: Three different vector fields on a square made of cables. The left hand one would result in a simple translation of the tensegrity. The center one indicates a rotation around the marked point. The right field is the only one which would change the shape of the tensegrity. The first two are elements of $T(p)$, the third is in $I(p)$ but not in $\bar{T}(p)$ or $T(p)$.

In the work we do here we will look for inspiration to Murray et al. (1994), who much of the same work, though primarily in $\mathbb{R}^3$.

Intuitively, a *rigid motion of space* is a continuous movement of $\mathbb{R}^n$ which preserves all distances and angles. More formally, H_t is a rigid motion of space if

$$H_t: [0, 1] \rightarrow \text{Aut } \mathbb{R}^n$$

(where $\text{Aut } \mathbb{R}^n$ is the set of automorphisms of $\mathbb{R}^n$) such that H_0 is the identity map on $\mathbb{R}^n$ and for all $v, w \in \mathbb{R}^n$ and $t \in [0, 1]$,

$$\langle H_t(v), H_t(w) \rangle = \langle v, w \rangle.$$

We will call a map $g: \mathbb{R}^n \rightarrow \mathbb{R}^n$ a *rigid transformation* if there is some rigid motion of space, H_t such that $g = H_1$.

Clearly, any such transformation must take an orthonormal basis for $\mathbb{R}^n$ to an orthonormal basis for $\mathbb{R}^n$, so we can describe it by telling where the origin goes and how the new set of coordinates is oriented with respect to the old. That is, it will consist of a translation and a “reorientation”. That reorientation may be a rotation (in fact, must be in $\mathbb{R}^2$ and $\mathbb{R}^3$) or a composition of rotations (Artin, 1991, p. 125).

We’ll look first at the reorientation and then see what it takes to include the translation. Taking as our original basis for $\mathbb{R}^n$ the traditional orthonormal basis consisting of unit vectors in the coordinate directions, $e_1, \dots, e_n$, we can describe a reorientation by taking the vectors $\hat{e}_1, \dots, \hat{e}_n$ of the new coordinate axes and building a reorientation matrix

$$R = [\hat{e}_1 \cdots \hat{e}_n].$$

Now since the $\hat{e}_1, \dots, \hat{e}_n$ form an orthonormal basis and hence are mutually orthogonal and of unit length, we have

$$R^\top R = RR^\top = \begin{bmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 1 \end{bmatrix} = I \text{ (the identity matrix).}$$

Then, since

$$1 = \det I = \det R^\top R = \det R^\top \det R = (\det R)^2, \quad (\text{A.1})$$

we have that $\det R = \pm 1$. Now if we think of the set of all $n \times n$ matrices as the space $\mathbb{R}^{n^2}$, then the determinant is a continuous function from that space into $\mathbb{R}$. But the set $\{-1, 1\}$ has two connected components, so the set of matrices described by Equation (A.1) must have at least two connected components.¹ Furthermore, since

¹It turns out to be exactly two (see Artin (1991, p. 124)), but we don’t need that information here.

the rigid motions provide a path between the identity transformation and the reorientation in question, all of the reorientations which play into our rigid transformations must lie in the connected component in which the identity has its home.

That means that we have $\det R = 1$ for every reorientation R . So our reorientations are contained in the set of “special orthogonal matrices”, which we can write formally as

$$SO_n = \{R \in \mathbb{R}^{n \times n} : R^\top R = I \text{ and } \det R = 1\}$$

(see, for example, [Artin \(1991, p. 124\)](#)).

On the other hand, if $R \in SO_n$, then for any $v, w \in \mathbb{R}^n$,

$$\langle Rv, Rw \rangle = v^\top R^\top R w = v^\top w = \langle v, w \rangle,$$

and we can produce a rigid motion of space by taking the rotations which make up R and varying their angles from 0 through the angle desired. Hence our set of reorientations is SO_n .

However, we’re not really interested in SO_n , *per se*. We want the “infinitesimal reorientations”, the directions in the space of $n \times n$ matrices in which a rigid motion of space can leave the origin – the tangent space to SO_n at the identity (more on that in a moment).

Now the $n \times n$ matrices form a normed linear space using the Frobenius norm (see, for example, [Artin \(1991, p. 153\)](#)), which is essentially the Euclidean norm with the matrices thought of as vectors in $\mathbb{R}^{n^2}$.

SO_n is a subset of that space of $n \times n$ matrices (not a subspace, since, for example, scaling a matrix changes its determinant), but it is more than that. It is a Lie group, which is to say, a group that is also a smooth manifold. So its tangent space at the identity is well-defined. This tangent space to SO_n at the identity is usually called $\mathfrak{so}_n$ ([Sharpe, 1997, pp. 12,64](#)).

Now the elements of a tangent space to a surface at a given point provide us the best linear approximation to that surface at that point. So if our surface is defined

as the level curve of some function $f(x) = c$, the elements of the tangent space are precisely those vectors v for which $f(x + \varepsilon v) - c$ is proportional to ε^2 rather than to ε .

In our case, SO_n is defined by the equation² $R^\top R = I$, so we can see it as the level curve $f(R) = 0$ for $f(R) = \|R^\top R - I\|$.

Then $dR \in \mathfrak{so}_n$ if and only if $\left. \frac{d}{d\varepsilon} f(I + \varepsilon dR) \right|_{\varepsilon=0} = 0$. That is,

$$\begin{aligned}
 0 &= \left. \frac{d}{d\varepsilon} f(I + \varepsilon dR) \right|_{\varepsilon=0} \\
 &= \left. \frac{d}{d\varepsilon} \|(I + \varepsilon dR)^\top (I + \varepsilon dR) - I\| \right|_{\varepsilon=0} \\
 &= \left. \frac{d}{d\varepsilon} \|I^\top I + \varepsilon dR^\top I + I^\top \varepsilon dR + \varepsilon^2 dR^\top dR - I\| \right|_{\varepsilon=0} \\
 &= \left. \frac{d}{d\varepsilon} \|\varepsilon dR^\top + \varepsilon dR + \varepsilon^2 dR^\top dR\| \right|_{\varepsilon=0} . \\
 &= \left. \frac{d}{d\varepsilon} \varepsilon \|dR^\top + dR + \varepsilon dR^\top dR\| \right|_{\varepsilon=0} . \\
 &= \|dR^\top + dR + \varepsilon dR^\top dR\| + \varepsilon \left. \frac{d}{d\varepsilon} \|dR^\top + dR + \varepsilon dR^\top dR\| \right|_{\varepsilon=0} .
 \end{aligned}$$

If we knew that $\left. \frac{d}{d\varepsilon} \|dR^\top + dR + \varepsilon dR^\top dR\| \right|_{\varepsilon=0}$ was finite, we'd know that that entire term is zero at $\varepsilon = 0$. Well, denoting the entry of dR in the i^{th} row and j^{th} column by d_{ij} and the equivalent entry in $dR^\top dR$ by c_{ij} gives us

$$\begin{aligned}
 &= \|dR^\top + dR + \varepsilon dR^\top dR\| + \varepsilon \left. \frac{\sum_{i=1}^n \sum_{j=1}^n (d_{ji} + d_{ij} + \varepsilon c_{ij}) c_{ij}}{\|dR^\top + dR + \varepsilon dR^\top dR\|} \right|_{\varepsilon=0} . \\
 &= \|dR^\top + dR\| \implies dR^\top + dR = \mathbf{0} \implies dR = -dR^\top .
 \end{aligned}$$

So $dR \in \mathfrak{so}_n$ if and only if dR is a skew-symmetric matrix. That means that the entries above the diagonal may be arbitrary but then the rest of the matrix is determined, so $\mathfrak{so}_n$ is $\frac{n(n-1)}{2}$ -dimensional. Since the dimension of a tangent space is the same as the dimension of the manifold (see, for example, [Kolář et al. \(1993, p.](#)

²There is also the determinant requirement, but that simply selects which connected component we are concerned with. As we are asking for the tangent space *at the identity*, we will get the same results if we think of just the component in hand or of both components together.

7)), that tells us also the dimension of SO_n . However, there's an intuitive argument for the same value.

Thinking back, we remember that the columns of any element of SO_n form an orthonormal basis for $\mathbb{R}^n$. So we start with n^2 degrees of freedom and apply n constraints to make the column vectors unit length. Then, for any pair of vectors, we add another constraint to make them orthogonal. That's $\binom{n}{2} = \frac{n(n-1)}{2}$ more constraints for a total of $\frac{n(n+1)}{2}$ degrees of freedom gone and $\frac{n(n-1)}{2}$ remaining.

A.2 THE SPECIAL EUCLIDEAN GROUP: SE_n AND $\mathfrak{se}_n$

Now that we've seen what the reorientations are like, we'd like to add in the translations. Of course, since any linear transformation on $\mathbb{R}^n$ takes the origin to the origin, translations are outside that realm. But by moving up one dimension, we can return to the world of linear transformations.

To do so we will take a point in $\mathbb{R}^n$, $p = [p_1 \ p_2 \ \cdots \ p_n]^\top$ and move it into $\mathbb{R}^{n+1}$ by adding a coordinate and setting its value to 1, thus: $[p_1 \ \cdots \ p_n \ 1]^\top$. So then, if $R \in SO_n$ tells how we wish to reorient the axes and $t \in \mathbb{R}^n$ tells where the origin should go, we can find the new location of the point p by the linear transformation:

$$\begin{bmatrix} \hat{p} \\ 1 \end{bmatrix} = \begin{bmatrix} R & t \\ \mathbf{0} & 1 \end{bmatrix} \begin{bmatrix} p \\ 1 \end{bmatrix} = \begin{bmatrix} Rp + t \\ 1 \end{bmatrix} \quad (\text{A.2})$$

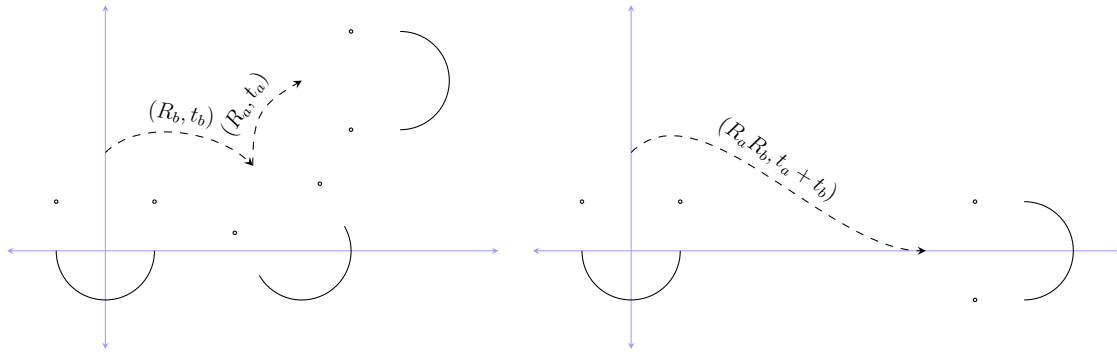
The matrices $\begin{bmatrix} R & t \\ \mathbf{0} & 1 \end{bmatrix}$ form the *Special Euclidean Group* SE_n (sometimes called just the *Euclidean Group*, Euc_n). If we label the $(n+1) \times (n+1)$ matrices in the fashion $\begin{bmatrix} R & t \\ c & d \end{bmatrix}$ (where R is an $n \times n$ matrix, t is an n -coordinate column vector, c an n -coordinate row vector and d a scalar), then SE_n is the subset of that space defined by $R^\top R = I$, $c = \mathbf{0}$ and $d = 1$.

Since we can pair any reorientation with any translation, SE_n appears to be isomorphic to $SO_n \times \mathbb{R}^n$ and so would be of dimension $\frac{n(n+1)}{2}$, and that's close to true. The dimension is right, but the group law is wrong.

Reorientations and translations don't commute. From Equation (A.2) on the preceding page we can see that an element of SE_n first applies the reorientation and then the translation, so suppose we have two elements of SE_n , say

$a = (R_a, t_a)$: “rotate around the origin by $\pi/3$ and then move right 1 unit”, and

$b = (R_b, t_b)$: “rotate around the origin by $\pi/6$ and then move right 1 unit”.



(a) The effect of applying (R_b, t_b) and then (R_a, t_a) .

(b) The effect of applying $(R_a R_b, t_a + t_b)$.

Figure A.2.1: SE_n is not isomorphic to $SO_n \times \mathbb{R}^n$, since the operations of rotating around the origin and translating do not commute.

Then, as Figure A.2.1 shows, the operations $(R_a, t_a)(R_b, t_b)$ and $(R_a R_b, t_a + t_b)$ are not, in general, the same. However, the reorientations are automorphisms of $\mathbb{R}^n$ and the translations are $\mathbb{R}^n$. So we can think of the reorientations acting on the translations and define a group operation thus:

$$(R_1, t_1)(R_2, t_2) = (R_1 R_2, t_1 + (R_1 t_2)).$$

The group thus formed is known as the *semidirect product* of SO_n and $\mathbb{R}^n$, denoted $SO_n \ltimes \mathbb{R}^n$ (see, for example, Rotman (1995, p. 167) or Gaal (1973, p. 236), for more information about semidirect products).

Proposition A.2.1. $SO_n \ltimes \mathbb{R}^n$ is a group and is isomorphic to SE_n .

Proof. Let $R_1, R_2, R_3 \in SO_n$ and $t_1, t_2, t_3 \in \mathbb{R}^n$. Then

$$\begin{aligned}
 (R_1, t_1)((R_2, t_2)(R_3, t_3)) &= (R_1, t_1)(R_2 R_3, t_2 + R_2 t_3) \\
 &= (R_1 R_2 R_3, t_1 + R_1 t_2 + R_1 R_2 t_3) \\
 &= (R_1 R_2, t_1 + R_1 t_2)(R_3, t_3) \\
 &= ((R_1, t_1)(R_2, t_2))(R_3, t_3)
 \end{aligned}$$

so we have associativity. The element $(I, \mathbf{0})$ is clearly the identity, so we only need inverses. Since SO_n is a group, we have inverses for all of its elements. Using these we can build the inverses we need:

$$\begin{aligned}
 (R_1, t_1)(R_1^{-1}, -R_1^{-1}t_1) &= (R_1 R_1^{-1}, t_1 + R_1(-R_1^{-1}t_1)) \\
 &= (I, t_1 - t_1) = (I, \mathbf{0}).
 \end{aligned}$$

So $SO_n \ltimes \mathbb{R}^n$ is a group. Consider the map $f: SO_n \ltimes \mathbb{R}^n \rightarrow SE_n$ by $(R, t) \mapsto \begin{bmatrix} R & t \\ \mathbf{0} & 1 \end{bmatrix}$.

Then

$$\begin{aligned}
 f((R_1, t_1))f((R_2, t_2)) &= \begin{bmatrix} R_1 & t_1 \\ \mathbf{0} & 1 \end{bmatrix} \begin{bmatrix} R_2 & t_2 \\ \mathbf{0} & 1 \end{bmatrix} \\
 &= \begin{bmatrix} R_1 R_2 & R_1 t_2 + t_1 \\ \mathbf{0} & 1 \end{bmatrix} = f((R_1, t_1)(R_2, t_2)).
 \end{aligned}$$

So the two groups are homomorphic. Furthermore, f is onto and has trivial kernel, by inspection, so $SO_n \ltimes \mathbb{R}^n$ and SE_n are isomorphic. $\square$

It turns out that SE_n is a Lie group, just as SO_n was (Sharpe, 1997, p. 64), which is good, as we are still interested in the tangent space, $\mathfrak{se}_n$. Since the elements of SE_n look like $\begin{bmatrix} R & t \\ c & d \end{bmatrix}$, the elements of $\mathfrak{se}_n$ look like $\begin{bmatrix} dR & dt \\ dc & dd \end{bmatrix}$. We already know that dR must be in $\mathfrak{so}_n$ and we can simply take differentials of the other two equations to give us $dc = \mathbf{0}$ and $dd = 0$. So the elements of $\mathfrak{se}_n$ look like $\begin{bmatrix} dR & dt \\ \mathbf{0} & 0 \end{bmatrix}$ where $dR \in \mathfrak{so}_n$ and $dt \in \mathbb{R}^n$.

All that remains is, then, is to figure out which vector fields correspond to the elements of $\mathfrak{se}_n$. To find those, we take the points in $p(\mathcal{V})$ and apply the elements of $\mathfrak{se}_n$ to them. In the next subsection, we give an example.

A.2.1 EXAMPLES

Consider the crossed-square example from [Figure 2.1.1](#) on page 15, shown also in [Figure A.2.2](#). We've shown that the elements of $\mathfrak{se}_2$ look like $\begin{bmatrix} 0 & a & x \\ -a & 0 & y \\ 0 & 0 & 0 \end{bmatrix}$. For the crossed square, $p(2) = (1, 0)$, so we multiply

$$\begin{bmatrix} 0 & a & x \\ -a & 0 & y \\ 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix} = \begin{bmatrix} x \\ -a + y \\ 0 \end{bmatrix}.$$

Thus we have $V(2) = (x, y - a)$. Following through the rest of the arithmetic gives us that the vector fields in $T(p)$ look like

$$V(v) = \begin{cases} (x, y), & v = 1 \\ (x, y - a), & v = 2 \\ (x + a, y - a), & v = 3 \\ (x + a, y), & v = 4. \end{cases}$$

For example, the two vector fields in [Figure A.2.2](#) (which are the same as the first

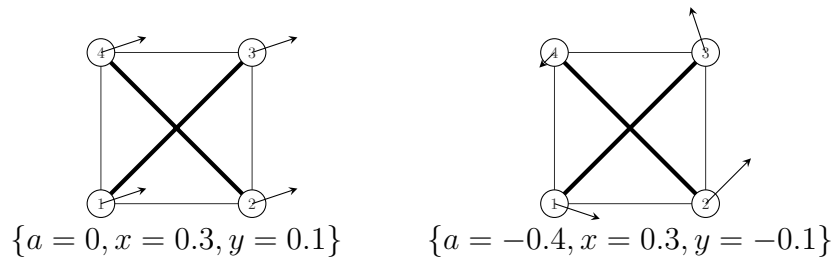


Figure A.2.2: Two vector fields in $T(p)$ on the crossed square.

two in [Figure A.1.1](#) on page 129) are given by $\{a = 0, x = 0.3, y = 0.1\}$ and $\{a = -0.4, x = 0.3, y = -0.1\}$ respectively.

APPENDIX B

MOTZKIN'S THEOREM

B.1 PREPARING TO USE MOTZKIN'S THEOREM

B.1.1 BACKGROUND

In [Section 3.1](#) on page 37, we talked about Theorems of the Alternative and how our main theorem relates to them. Here we will use one of them to give an alternate proof for the first main theorem. The new version actually has stronger hypotheses, so it's not as strong a theorem, but we include it here because the technique is interesting.

The theorems of Stiemke ([3.1.1](#)) and Gordan ([3.1.2](#)) in [Section 3.1](#) are for finite-dimensional Euclidean space. We need them extended to more general spaces. Unfortunately, a literature search failed to locate generalized versions.

However, Mangasarian ([1969](#)) proves Stiemke's theorem using a theorem from Theodore Motzkin's dissertation ([1936](#)) commonly known as Motzkin's Transposition Theorem. And B. D. Craven ([1978](#)) gives a version of Motzkin's Theorem for normed linear spaces.

The theorem reads as follows (here $L(X, W)$ denotes the space of continuous linear maps from X to W and A^T is Craven's notation for the adjoint of A):

Theorem B.1.1 (Motzkin Transposition Theorem). *Let X, W, Z be normed spaces; let $S \subset W$ be a convex cone, with $\text{int } S \neq \emptyset$; let $T \subset Z$ be a closed convex cone; let $A \in L(X, Z)$ and $B \in L(X, W)$. If the convex cone $A^T(T^*)$ is weak-* closed, then exactly one of the two following systems has a solution:*

$$(I) \quad -Ax \in T, \quad -Bx \in \text{int } S \quad (x \in X);$$

(II) $p \circ B + q \circ A = 0, q \in T^*, 0 \neq p \in S^*$.

Proof. See Craven (1978, p. 32). □

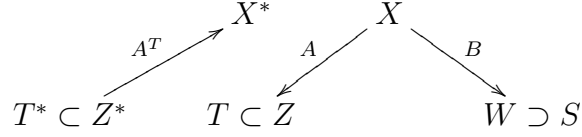


Figure B.1.1: The diagram for the Motzkin Transposition Theorem.

Figure B.1.1 shows the diagram for Motzkin's theorem. We know that $C(\mathcal{E})^+$ is a convex cone with nonempty interior, so it seems reasonable to set $S = C(\mathcal{E})^+$ and, of course, $W = C(\mathcal{E})$. In that case, if we made B be the negative of our rigidity operator and set $X = \text{VF}(\mathcal{V})$, then $-Bx \in \text{int } S$ would mean $YV \in \text{int } C(\mathcal{E})^+$, that is, V is a strictly positive motion.

That seems promising, but what are T and Z ? Well, that $-Ax \in T$ could allow us to implement the design variations. If $T = \mathcal{X}$ (which, of course, means that $Z \in \text{VF}(\mathcal{V})$) and $A = -\text{Id}$, the negative of the identity map, then $-Ax \in T$ becomes $V \in \mathcal{X}$.

How does that work out with the other alternative? Using the definitions we have put forth, we get Case (II) as $-pY - q = 0$, where $0 \neq p \in C^*(\mathcal{E})^+$ and $q \in \mathcal{X}^*$. So we could take p as our stress μ and (so long as $\mathcal{X}$ is a subspace so that $\mathcal{X}^* = \mathcal{X}^\perp$), we get the statement $Y^*\mu \in \mathcal{X}^\perp$. That is, there exists a semipositive stress. So this seems a good match for our needs. Our only concerns are about having $\mathcal{X}^\perp$ be closed and the weak-* closure of T^* . To make those easier to answer, we'll make $\text{VF}(\mathcal{V})$ into a Hilbert space.

B.1.2 γ AND $\text{VF}(\gamma)$

In what follows, we will work exclusively with tensegrities whose vertices lie along a simple, rectifiable (in fact, arclength parametrized) curve (again called γ) with finitely many connected components and total length ℓ .

We want $\text{VF}(\gamma)$ to be a Hilbert space, so we need to endow it with an inner product. In doing this, we'll follow the lead of John Pardon ([private communication](#)). We'll redefine $\text{VF}(\gamma)$ to be a subset of the absolutely continuous functions and establish the inner product $\langle V_1, V_2 \rangle = \int V_1' \cdot V_2' ds$ on it.

For that inner product to be finite, we'll need $\int |V'|^2 ds < \infty$ for all $V \in \text{VF}(\gamma)$, and that to get completeness, we'll want the derivative map on $\text{VF}(\gamma)$ to be 1-1.

Pardon accomplishes that last task by requiring that $V(0) = \mathbf{0}$, but we take a different route. We would like $\int_0^\ell V(s) ds = \mathbf{0}$, that is, we want the “center of mass” of the variation to be at $\mathbf{0}$. This will have the effect that $V(s)$ as a variation will give γ no net translation. In summary:

Definition B.1.2. The space of variations on γ , $\text{VF}(\gamma)$ will consist of all absolutely continuous vector fields V on γ which satisfy $\int |V'|^2 < \infty$ and $\int V = \mathbf{0}$.

It may help to have a more explicit definition for elements of $\text{VF}(\gamma)$.

Proposition B.1.3. $\text{VF}(\gamma)$, as defined in Definition [B.1.2](#) consists of those vector fields on γ which satisfy $\int |V'|^2 < \infty$ and

$$V(s) = \int_0^s V'(t) dt - \frac{1}{\ell} \int_0^\ell \int_0^t V'(r) dr dt. \quad (\text{B.1})$$

Proof. By Theorem [3.7.17](#) on page [82](#), vector fields which satisfy Equation (B.1) are absolutely continuous, so we need only show that this is the equation which gives us $\int V = \mathbf{0}$. We know that $V(s) = \int_0^s V'(t) dt + V(0)$. So

$$\int_0^\ell V(s) ds = \int_0^\ell \int_0^s V'(t) dt ds + \int_0^\ell V(0) ds = \int_0^\ell \int_0^s V'(t) dt ds + \ell V(0)$$

To get $\int_0^\ell V(s) ds = 0$, then, we need

$$V(0) = -\frac{1}{\ell} \int_0^\ell \int_0^s V'(t) dt ds.$$

□

Now we are ready to establish the nature of our new $\mathbf{VF}(\gamma)$.

Proposition B.1.4. $\mathbf{VF}(\gamma)$ is a Hilbert space under the inner product $\langle V_1, V_2 \rangle = \int V'_1 \cdot V'_2 ds$.

Proof. First we need to establish that $\langle V_1, V_2 \rangle$ as we've defined it really is an inner product. We have three things to check:

(i) Symmetry:

$$\langle V_1, V_2 \rangle = \int V'_1 \cdot V'_2 ds = \int V'_2 \cdot V'_1 ds = \langle V_2, V_1 \rangle$$

(ii) Positive definiteness:

$$\langle V, V \rangle = \int V' \cdot V' ds \geq 0$$

We note that $\langle V, V \rangle = 0$ only when $V' = \mathbf{0}$ almost everywhere. So V must be constant. But because $\int V = \mathbf{0}$, this only happens when $V = \mathbf{0}$. Also, by our definition of $\mathbf{VF}(\gamma)$, $\int \|V'\|^2 < \infty$, so $\langle V, V \rangle$ is finite.

(iii) Linearity in the first position:

$$\begin{aligned} \langle aV_1 + bV_2, V_3 \rangle &= \int (aV'_1 + bV'_2) \cdot V'_3 ds \\ &= a \int V'_1 \cdot V'_3 ds + b \int V'_2 \cdot V'_3 ds = a\langle V_1, V_3 \rangle + b\langle V_2, V_3 \rangle \end{aligned}$$

(for all $a, b \in \mathbb{R}$ and $V_1, V_2, V_3 \in \mathbf{VF}(\gamma)$).

So $\langle V_1, V_2 \rangle$ is an inner product.

Next we need to show that our space is complete with the norm $\|V\| = \langle V, V \rangle^{\frac{1}{2}}$.

Suppose that there is a Cauchy sequence $\{V_n\}$ of elements of $\text{VF}(\gamma)$. That is, for any $\varepsilon > 0$, there exists some positive integer N such that whenever $m, n > N$, we have

$$\varepsilon > \|V_m - V_n\| = \left(\int \|V'_m - V'_n\|^2 \right)^{1/2}.$$

But we recognize that final term as $\|V'_m - V'_n\|_2$, the L^2 norm of $V'_m - V'_n$. L^2 is complete under its norm (see, for example, [Folland \(1999, p. 183\)](#)), so there is some $V' \in L^2$ such that $V'_n \rightarrow V'$. With that in hand we can calculate $V(s)$ from [Equation \(B.1\)](#) on page 139. The result is an absolutely continuous (by Theorem 3.7.17 on page 82) vector field for which $\int |V'_n|^2 < \infty$ and $\int V = 0$. Thus V is an element of $\text{VF}(\gamma)$. And so $\text{VF}(\gamma)$ with our inner product is a Hilbert space. $\square$

One delightful aspect of Hilbert spaces is that every linear functional on a Hilbert space is given by inner product with some element of the space. This is the result of another theorem also called the “Riesz Representation Theorem”, which we give below. This allows us to think of elements of a Hilbert space and its dual as lying in the same space.

Theorem B.1.5 (Riesz Representation Theorem for Hilbert Spaces). *Let X be a Hilbert space. Then, for every $f \in X^*$, there exists a unique element $p \in X$ such that $f(x) = \langle x, p \rangle$ for all $x \in X$ and $\|f\| = \|p\|$.*

Proof. See [Ponnusamy \(2002, p. 430\)](#). $\square$

B.1.3 $\mathcal{X}$ AND $\mathcal{X}^\perp$

Having established a new $\text{VF}(\gamma)$, we choose our “design variations” $\mathcal{X}$ as follows:

$$\mathcal{X} = \{V \in \text{VF}(\gamma) : V'(s) \cdot \gamma'(s) = 0 \text{ almost everywhere} \}. \quad (\text{B.2})$$

We will need to check that $\mathcal{X}$ is a subspace of our new $\text{VF}(\gamma)$. Furthermore, we’ll need $\mathcal{X}$ to be closed. In proving that it is, we’ll make use of the Cauchy-Schwarz inequality:

Theorem B.1.6 (Cauchy-Schwarz Inequality). *For all x, y in a Hilbert space, $\langle x, y \rangle \leq \|x\| \|y\|$.*

Proof. See Royden (1968, p. 210). □

Lemma B.1.7. $\mathcal{X}$, as defined in Equation (B.2), is a closed subspace of $\mathbf{VF}(\gamma)$.

Proof. First we'll dispose of the subspace portion. Let $V, W \in \mathcal{X}$ and $\alpha, \beta \in \mathbb{R}$. Then

$$(\alpha V + \beta W)'(s) \cdot \gamma'(s) = \alpha V'(s) \cdot \gamma'(s) + \beta W'(s) \cdot \gamma'(s) = 0$$

almost everywhere, putting $\alpha V + \beta W \in \mathcal{X}$. Now for closure.

Let $\{V_n\} \in \mathcal{X}$ be a sequence of vector fields that limits to some vector field $V \in \mathbf{VF}(\gamma)$. Then

$$\begin{aligned} \int (V'(s) \cdot \gamma'(s))^2 ds &= \int (V'(s) \cdot \gamma'(s) - V_n'(s) \cdot \gamma'(s))^2 ds \text{ since } V_n \in \mathcal{X} \\ &= \int ((V'(s) - V_n'(s)) \cdot \gamma'(s))^2 ds \\ &\leq \int \|V'(s) - V_n'(s)\|^2 \|\gamma'(s)\|^2 ds \text{ by Cauchy-Schwarz} \\ &= \int \|V'(s) - V_n'(s)\|^2 ds \text{ since } \gamma \text{ is unit speed} \\ &= \|V - V_n\|^2 \end{aligned}$$

But $V_n \rightarrow V$, so for any $\varepsilon > 0$, there is a positive integer N such that $\|V - V_n\| < \sqrt{\varepsilon} \Rightarrow \|V - V_n\|^2 < \varepsilon$. That means that $\int (V'(s) \cdot \gamma'(s))^2 ds$ is nonnegative, but less than every positive number. Hence,

$$\int (V'(s) \cdot \gamma'(s))^2 ds = 0$$

and thus $V'(s) \cdot \gamma'(s) = 0$ almost everywhere. So $V \in \mathcal{X}$, and $\mathcal{X}$ is closed. □

Lemma B.1.8. *The dual cone of a subspace X is its annihilator. That is, if X is a subspace, $X^* = X^\perp$.*

Proof. Let $\mu \in X^*$. Then $\mu \cdot x \geq 0$ by definition. However, since X is a subspace, $x \in X \Rightarrow -x \in X$, so $\mu \cdot (-x) = -\mu \cdot x$ must also be nonnegative. But that means that $\mu \cdot x = 0$. That is, $\mu \in X^\perp$.

Conversely, let $\eta \in X^\perp$. Then, for all $x \in X$, we have $\eta \cdot x = 0 \geq 0$, so $\eta \in X^*$. $\square$

Lemma B.1.9. *Since $\mathcal{X}$ is a subspace, $\mathcal{X}^\perp$ and $Y(\mathcal{X})$ are subspaces.*

Proof. Let $\alpha, \beta \in \mathbb{R}$ and $V_1, V_2 \in \mathcal{X}^\perp$. Then, for all $V \in \mathcal{X}$, we have

$$(\alpha V_1 + \beta V_2) \cdot V = \alpha V_1 \cdot V + \beta V_2 \cdot V = 0,$$

so $\mathcal{X}^\perp$ is a subspace.

Similarly, if $YV, YW \in Y(\mathcal{X})$, then $\alpha V + \beta W \in \mathcal{X}$ (since $\mathcal{X}$ is a subspace) and $\alpha YV + \beta YW = Y(\alpha V + \beta W)$. So $Y(\mathcal{X})$ is a subspace. $\square$

B.2 MOTZKIN'S THEOREM APPLIED

We are now ready to apply the Motzkin Transposition Theorem to our setup. For clarity, [Figure B.2.1](#) shows the Motzkin Diagram and our situation given as a parallel diagrams.

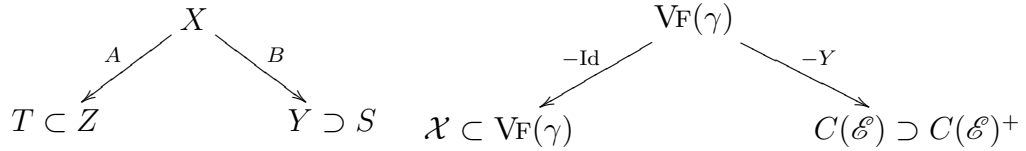


Figure B.2.1: The Motzkin diagram and our matching situation.

We need to understand what “ $A^T(T^*)$ is weak-* closed” means in our situation. For us, A and A^T are both negative identity functions, one on $\text{VF}(\gamma)$ and the other on $\text{VF}^*(\gamma)$, which, thanks to the Riesz Representation Theorem for Hilbert Spaces (Theorem B.1.5 on page 141), we can identify with $\text{VF}(\gamma)$.

We showed, in Lemma B.1.8, that $\mathcal{X}^* = \mathcal{X}^\perp$, so $A^T(T^*)$ is simply $-\mathcal{X}^\perp$, which (by Lemma B.1.9) is just $\mathcal{X}^\perp$. Now for the theorem.

Theorem B.2.1. *A tensegrity $G(p)$ is partially $\mathcal{X}$ -bar equivalent if and only if $G(p)$ has a semipositive stress.*

Proof. $\text{VF}(\gamma)$ and $C(\mathcal{E})$ are normed spaces, $C(\mathcal{E})$ with the sup norm and $\text{VF}(\gamma)$ with the norm arising from its inner product. So we have the right kind of spaces for Motzkin's Theorem (B.1.1 on page 137). For its hypotheses, we need the following to be true:

- (a) $C(\mathcal{E})^+$ is a convex cone with nonempty interior,
- (b) $\mathcal{X}$ is a closed convex cone and
- (c) $\mathcal{X}^\perp$ is closed.

If those are true, then exactly one of the follow systems has a solution:

- (I) $V \in \mathcal{X}, YV \in \text{int } C(\mathcal{E})^+$;
- (II) $Y^*\mu = V, -V \in \mathcal{X}^\perp, \mathbf{0} \neq \mu \in C^*(\mathcal{E})^+$.

That is, $G(p)$ has either a strictly positive motion (I) or a semipositive stress (II).

$\mathcal{X}$ is not only a closed convex cone, but moreover a closed subspace, by Lemma B.1.7, so Item (b) is satisfied. Since $\text{VF}(\gamma)$ is a Hilbert space, the subspace $\mathcal{X}^\perp$ (since it is defined by orthogonality to a set) is closed (see, for example, Folland (1999, p. 173)). That takes care of Item (c).

It remains only to show that $C(\mathcal{E})^+$ is a convex cone with nonempty interior. Certainly, if f is a nonnegative continuous function from $\mathcal{E}$ to $\mathbb{R}$, then αf is as well, for any $\alpha > 0$. Likewise, if $f, g \in C(\mathcal{E})^+$, then any convex combination of f and g is also in $C(\mathcal{E})^+$, so $C(\mathcal{E})^+$ is a convex cone. Furthermore, by Lemma 3.3.2 on page 53, $\text{int } C(\mathcal{E})^+$ contains all of the strictly positive functions (and thus is nonempty).

By applying Motzkin's theorem, we have our result. □

B.3 MOTZKIN ON THE OTHER HAND

Having had such success with that direction, it seems only natural to see if we can get anything more out of Motzkin's Theorem. What we've proven is

$$\text{there is no strictly positive motion} \Leftrightarrow \text{there is a semipositive stress.}$$

Can we use the Motzkin Transposition Theorem to show that

$$\text{there is no semipositive motion} \Leftrightarrow \text{there is a strictly positive stress?}$$

The natural thing would be to try to require that our measures fall in the interior of $C^*(\mathcal{E})^+$. But we saw in Theorem 3.6.3 on page 71 that when $\mathcal{E}$ is infinite, $C^*(\mathcal{E})^+$ has no interior. So any success along these lines must be limited to the finite case. Let's give it a try.

We need $\text{int } S$ to be $\text{int } C^*(\mathcal{E})^+$, but in our basic setup, Y^* maps from $C^*(\mathcal{E})^+$ and nothing maps to it. So that's where we'll send our identity function (see Figure B.3.1).

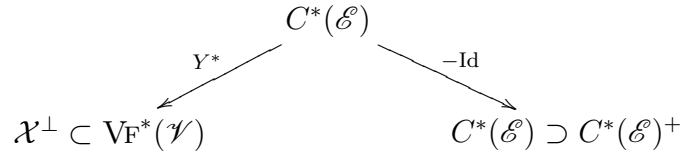


Figure B.3.1: Our new use of the Motzkin theorem.

Now, we need $C^*(\mathcal{E})^+$ to have non-empty interior, which it does. We need $\mathcal{X}^\perp = \{\mathbf{0}\}$ to be a closed convex cone, which it is. Finally, we need

$$Y^{**}((\mathcal{X}^\perp)^*) = Y(\text{VF}(\mathcal{V}))$$

to be weak-* closed. But Y is a finite-dimensional linear map and so its image is a subspace of the finite-dimensional $C(\mathcal{E})$ and thus it is closed.

Now we can apply Motzkin's theorem and get that exactly one of these two systems has a solution:

$$(I) \quad -Y^*\mu = \mathbf{0}, \mu \in \text{int } C^*(\mathcal{E})^+$$

$$(II) \quad YV = f, V \in \text{VF}(\mathcal{V}), 0 \neq f \in C(\mathcal{E})^+.$$

That is, “either there is a strictly positive stress or else there is a semipositive motion”.

Thus we have a new proof of Roth and Whiteley’s theorem.

B.4 STRESSES AS VARIATIONS

One somewhat surprising effect of making $\text{VF}(\gamma)$ into a Hilbert space is that any given stress μ has a variation, $V_\mu \in \mathcal{X}^\perp$ associated with it. That is because $Y^*\mu \in \mathcal{X}^\perp \subset \text{VF}^*(\gamma)$ and $\text{VF}(\gamma)$ and $\text{VF}^*(\gamma)$ can be identified with each other, thanks to the Riesz Representation Theorem for Hilbert Spaces (Theorem B.1.5 on page 141).

For example, in the case of the circle of struts, we found a stress $d\theta$, which was uniform on the struts. $Y^*d\theta$, then corresponds to some variation $V_{d\theta}$ in such way that

$$\begin{aligned} \langle V_{d\theta}, V \rangle &= (Y^*d\theta)V = d\theta(YV) = \int_0^\pi YV d\theta \\ &= \int_0^\pi (V(\theta) - V(\theta + \pi)) \cdot ((\cos \theta, \sin \theta) - (\cos(\theta + \pi), \sin(\theta + \pi))) d\theta \\ &= \int_0^\pi (V(\theta) - V(\theta + \pi)) \cdot 2(\cos \theta, \sin \theta) d\theta \\ &= \int_0^\pi V(\theta) \cdot 2(\cos \theta, \sin \theta) d\theta - \int_0^\pi V(\theta + \pi) \cdot 2(\cos \theta, \sin \theta) d\theta \\ &= \int_0^\pi V(\theta) \cdot 2(\cos \theta, \sin \theta) d\theta - \int_\pi^{2\pi} V(\theta) \cdot 2(\cos(\theta - \pi), \sin(\theta - \pi)) d\theta \\ &= \int_0^\pi V(\theta) \cdot 2(\cos \theta, \sin \theta) d\theta + \int_\pi^{2\pi} V(\theta) \cdot 2(\cos \theta, \sin \theta) d\theta \\ &= \int_0^{2\pi} V(\theta) \cdot 2(\cos \theta, \sin \theta) d\theta. \end{aligned} \tag{B.3}$$

On the other hand, our inner product on $\text{VF}(\gamma)$ says that

$$\langle V_{d\theta}, V \rangle = \int_0^{2\pi} V'_{d\theta}(s) \cdot V'(s) ds.$$

We can integrate this by parts to get

$$\langle V_{d\theta}, V \rangle = V'_{d\theta}(s) \cdot V(s) \Big|_0^{2\pi} - \int_0^{2\pi} V''_{d\theta}(s) \cdot V(s) ds. \quad (\text{B.4})$$

Now, 0 and 2π are the same point on our domain, so that first term is zero. We can combine Equation (B.3) and Equation (B.4) to get

$$- \int_0^{2\pi} V''_{d\theta}(s) \cdot V(s) ds = \int_0^{2\pi} V(\theta) \cdot 2(\cos \theta, \sin \theta) d\theta,$$

which must be true for all $V \in \text{VF}(\gamma)$. Since $s = \theta$ on the unit circle, that gives us that $-V''_{d\theta}(\theta) = 2(\cos \theta, \sin \theta)$ or

$$V_{d\theta}(\theta) = 2(\cos \theta, \sin \theta),$$

which is shown in Figure B.4.1.

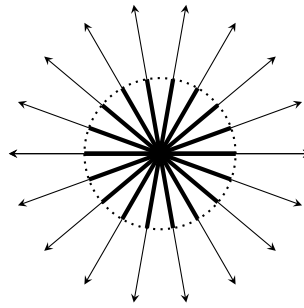


Figure B.4.1: The circle of struts with the vectors of $V_{d\theta}$.

Thinking back to our original definition of stress, this might seem the obvious thing. After all, it just looks to be the weighted edge vectors at each vertex. However, on second glance it seems less obvious. After all, this example has a high degree of symmetry. We were able to transform an integral on the edges into an integral on the vertices in the form $\int V \cdot (\text{something})$ which we could then relate to our inner product.

This is an area which bears further investigation, both in identifying the variation for a given stress and in understanding the interplay between that and the “curvature force” of Cantarella et al. (2006).

INDEX

~Symbols~

$\langle \cdot, \cdot \rangle$ (inner product)	14
$v \cdot w$ (dot product)	14
$X \setminus Y$ (set-theoretic difference)	14
$\overline{\quad}$ (bar)	15
--- (cable)	15
— (strut)	15
$\circ$ (vertex)	15
$\cdots$ (vertices)	15
$\cup$ (union)	16
$\sqcup$ (disjoint union)	16
$S^\perp$ (annihilator of S)	22, 31
S^* (dual cone of S)	22, 31
$\lfloor x \rfloor$ (greatest integer $\leq x$)	72

~ $\mathcal{A}$ ~

absolutely continuous	76
affine	
hull	27
set	27
algebraic dual	22
ambient space	
increasing dimension	116
annihilator	22, 31
Architecture	2
arclength parameterization	75
Augustin Louis Cauchy	5
Axiom of Choice	90

~ $\mathcal{B}$ ~

B (bars)	3, 23
$\overline{B}$ (bars in $\overline{G(p)}$)	23
bar	3
as strut-cable pair	16
bar equivalent	8, 20, 34
$\Leftrightarrow$ strictly positive stress	5, 9, 23, 35
with respect to $\mathcal{X}$.. see $\mathcal{X}$ -bar equivalent	
bar framework	3
Ben Roth	3, 14–36
Biology, Cellular	2
bone, human	1
Branko Grünbaum	6
brick	1
Bruce D. Craven	137
Buckminster Fuller	2

~ $\mathcal{C}$ ~

C (cables)	3, 16, 23
$\mathcal{C}$ (functional cables, $\mathbf{C} \cup \mathbf{B}$)	16, 23
$C(\mathcal{E})$ (cont. func. on $\mathcal{E}$)	18, 23, 52–54
$C(\mathcal{E})^+$ and $C(\mathcal{E})^-$ (orthants of $C(\mathcal{E})$)	18, 23

interior	53
$C^*(\mathcal{E})$ (top. dual of $C(\mathcal{E})$)	21, 23, 52–60
$C^*(\mathcal{E})^+$ (nonneg. orth. of $C^*(\mathcal{E})$)	22, 23
cable	3
Cantarella, Jason	7
Carpenter's Rule Problem	5
Cauchy, Augustin Louis	5
Cauchy-Schwarz Inequality	142
Cellular Biology	2
circle of struts .. 10, 73, 82–83, 85, 87, 91, 147	
compression	1–4, 8, 21, 23, 99
concrete	1
prestressed	4, 8, 21
cone, dual	22, 31, 142
Connelly, Robert	5–6
continuous	76
absolutely	76
Lipschitz	76
uniformly	76
continuous tensegrity	3
convex	
hull	29
set	27
finitely generated	30
countably covered	11, 90, 91, 121
$\Rightarrow$ strictly positive stress	11, 90
Craven, Bruce D.	137
crossed square	15–24, 24–26, 47, 70, 89, 108–110, 136
curvature force	147

~ $\mathcal{D}$ ~

Dairy Science	2
Dance	2
$d_{\mathcal{E}}$ (metric on $\mathcal{E}$)	43
Demaine, Erik D.	5
design variation	8, 18
direct product	39
direction	29
discrete topology	17, 52
Dominated Convergence Theorem	79, 80
dual	
algebraic	22
topological	22
dual cone	22, 31, 142
$d_{\mathcal{V}}$ (metric on $\mathcal{V}$)	39
$d_{\mathcal{V} \times \mathcal{V}}$ (metric on $\mathcal{V} \times \mathcal{V}$)	39
$d_{(\mathcal{V} \times \mathcal{V})/\sim}$ (metric on $(\mathcal{V} \times \mathcal{V})/\sim$)	40

~ $\mathcal{E}$ ~

$\mathcal{E}$ (edgeset)	9, 16, 23, 43, 39–52
-------------------------------	----------------------

Hausdorff.....	41, 42
non-compact	49–50
$\mathcal{E}_c(\mathcal{S} \sqcup \mathcal{E})$	51
edgeset	<i>see</i> $\mathcal{E}$
Erich Stiemke.....	38
Erik D. Demaine.....	5
Euclidean	
group.....	133
motion.....	17, 94, 95, 129–136
norm	52, 53
extended metric.....	43

~ $\mathcal{F}$ ~

face (of a convex set)	28
finitely generated convex set	30
force, curvature.....	147
framework, bar.....	3
Fu, Joseph H. G.....	7
Fuller, Buckminster	2
functional struts and cables.....	16

~ $\mathcal{G}$ ~

$G(p)$ (tensegrity).....	16, 23
$\overline{G}(p)$ ($G(p)$'s bar framework).....	16, 23
γ (synonym for $\mathcal{V}$)	24
Geoffrey C. Shephard	6
Gordan, Paul	38
granite.....	1
group	
Lie	131, 135
special Euclidean	133
special orthogonal	131
Grünbaum polygons	6
Grünbaum, Branko	6
Günter Rote	5

~ $\mathcal{H}$ ~

Hahn-Banach Theorem	64, 66
hull	
affine.....	27
convex	29
human bone	1

~ I ~

$I(p)$ (motions).....	18, 23
$\neq \overline{I}(p)$	20
$\supset \overline{I}(p) \supset T(p)$	19
$\overline{I}(p)$ (motions of $\overline{G}(p)$).....	19, 24
$\neq I(p)$	20
$\neq T(p)$	20
$\subset I(p)$	19
$\supset T(p)$	19
infinitesimal motion	<i>see</i> motion
infinitesimal rigidity	34
infinitesimally rigid.....	6, 8, 20, 95
with respect to $\mathcal{X}$	<i>see</i> $\mathcal{X}$ -infinitesimally rigid
interior	
of $C^*(\mathcal{E})^+$	71–73

of $C(\mathcal{E})^+$ and $C(\mathcal{E})^-$	53
relative.....	27

~ $\mathcal{J}$ ~

Jason Cantarella.....	7
John M. Sullivan.....	7
John Pardon.....	5
Joseph H. G. Fu	7

~ $\mathcal{K}$ ~

Kenneth D. Snelson	2, 13
key lemma (Roth & Whiteley).....	34–35
Kusner, Robert B.	7

~ $\mathcal{L}$ ~

length (of a curve).....	75
Lie group	131, 135
Lipschitz constant	76
Lipschitz continuous.....	76
load.....	8, 18, 19
semipositive.....	19
strictly positive	19

~ $\mathcal{M}$ ~

Mangasarian, Olvi	38, 137
maple, sugar	1
measure	
nonnegative	55
positive	55
regular.....	54
semipositive.....	55
strictly positive	55
metallurgy	2
metric	
extended	43
on $\mathcal{E}$	43
on $\mathcal{V}$	39
on $\mathcal{V} \times \mathcal{V}$	39
on $(\mathcal{V} \times \mathcal{V})/\sim$	40
on $C(\mathcal{E})$	53
on $C^*(\mathcal{E})$	53
metric topology.....	42, 43, 53, 72
minimally $\mathcal{X}$ -bar equivalent	86
$\Rightarrow$ strictly positive stress.....	10, 86
motion.....	8, 18
Euclidean	17, 94, 95, 129–136
rigid, of $\mathbb{R}^n$	<i>see</i> Euclidean motion
rigid, of space	8, 129
semipositive	8, 18
strictly positive.....	8, 18
Motzkin Transposition Theorem	9, 137
Motzkin, Theodore S.	9, 137
musculoskeletal tensegrity.....	2

~ $\mathcal{N}$ ~

Nancy Wrinkle.....	7
non-compact edgeset	49–50
nonnegative	
measure.....	55

vector	14
vector field	15
norm	
Euclidean	52, 53
on $C(\mathcal{E})$	52
on $C^*(\mathcal{E})$	52
operator	52
sup	52, 53, 60

$\sim \mathcal{O} \sim$

oak, pin	1
Olvi Mangasarian	38, 137
1 (constant function 1)	60, 66
operator norm	52
Ornithology	2

$\sim \mathcal{P} \sim$

p (map from $\mathcal{V}$ to $\mathbb{R}^n$)	3, 24
parameterization, arclength	75
Pardon, John	5
partially bar equivalent	8
partially $\mathcal{X}$ -bar equivalent	8, 21
$\Leftrightarrow$ semipositive stress	5, 9, 65, 144
patent, tensegrity	2
Paul Gordan	38
polyhedral set	28
positive measure	55
prestressed concrete	4, 8, 21
product	
direct	39
semidirect	134

$\sim \mathcal{R} \sim$

R. Tyrrell Rockafellar	26–30
rectifiable	75
regular measure	54
relative interior	27
Riesz Representation Theorem	55
for Hilbert Spaces	141, 143, 146
rigid	
infinitesimally	8, 20, 95
motion of $\mathbb{R}^n$	see Euclidean motion
motion of space	8, 129
transformation	130
$\mathcal{X}$ -infinitesimally	8
rigidity operator	18, 60
rigidity, infinitesimal	34
$\mathbb{R}^n$ (Euclidean n -space)	14
Robert B. Kusner	7
Robert Connelly	5–6
Robot Kinematics	2
Rockafellar, R. Tyrrell	26–30
Rote, Günter	5
Roth, Ben	3, 14–36

$\sim \mathcal{S} \sim$

S (struts)	3, 16, 23
$\mathcal{S}$ (functional struts, $\mathbf{S} \cup \mathbf{B}$)	16, 23
semidirect product	134
semipositive	

load	19
measure	55
motion	8, 18
stress	8, 22
$\Leftrightarrow$ partially $\mathcal{X}$ -bar equiv.	5, 9, 65, 144
vector	15
vector field	15
SE_n (Special Euclidean Group)	133
$\mathfrak{se}_n$ (tangent space to SE_n)	135
set	

affine	27
convex	27
edge	see $\mathcal{E}$
finitely generated convex	30
polyhedral	28
vertex	see $\mathcal{V}$
Shephard, Geoffrey C.	6
sleep disorders	2
Snelson, Kenneth D.	2, 13
SO_n (Special Orthogonal Group)	131
$\mathfrak{so}_n$ (tangent space to SO_n)	131
special Euclidean group	133
special orthogonal group	131
steel	1, 2
Stiemke, Erich	38
stone	1
stress	4, 8, 21, 22
semipositive	8, 22
$\Leftrightarrow$ partially $\mathcal{X}$ -bar equiv.	5, 9, 65, 144
strictly positive	8, 22
$\Leftarrow$ countably covered	11, 90
$\Leftarrow$ minimally $\mathcal{X}$ -bar equiv.	10, 86
$\Leftrightarrow$ bar equivalent	5, 9, 35
$\Rightarrow$ $\mathcal{X}$ -bar equivalent	9, 68

strictly positive	
load	19
measure	55
motion	8, 18
stress	8, 22
$\Leftarrow$ countably covered	11, 90
$\Leftarrow$ minimally $\mathcal{X}$ -bar equiv.	10, 86
$\Leftrightarrow$ bar equivalent	5, 9, 23, 35
$\Rightarrow$ $\mathcal{X}$ -bar equivalent	9, 68
vector	14
vector field	15
strut	3
Sullivan, John M.	7
sup norm	52, 53, 60
support	
of a function	53
of a measure	56
is closed	56

$\sim \mathcal{T} \sim$

$T(p)$ (Euclidean motions)	17, 24
$\neq \bar{I}(p)$	20
$\subset \bar{I}(p) \subset I(p)$	19
$\subset \mathcal{X}$	18
$\mathcal{T}_{\mathcal{E}}$ (topology on $\mathcal{E}$)	44

$\mathcal{T}_{\mathcal{E}}$ (topology on $\mathcal{E}$)	45
tensegrity	3
continuous	3
musculoskeletal	2
origin of term	2
tensegrity patent	2
tension	1–4, 8, 21, 23, 99
tensional integrity	2
Theodore S. Motzkin	9, 137
Theorem	
Dominated Convergence	79, 80
Gordan	38, 137
Hahn-Banach	64, 66
Motzkin Transposition	9, 137
of the Alternative	9, 35, 37–38, 121, 137
Riesz Representation	55, 66
for Hilbert Spaces	141, 143, 146
Schur	6
Stiemke	38, 137
topological dual	22
topology	
discrete	17, 52
metric	42, 43, 53, 72
$\mathcal{T}_{p(\mathcal{V})}$ (alt. topology on $\mathcal{V}$)	44
transformation, rigid	130
$\mathcal{T}_{sqp}$ (alt. topology on $\mathcal{E}$)	45
$\mathcal{T}_{\mathcal{S}}$ (topology on $\mathcal{S}$)	44
$\mathcal{T}_{\mathcal{V}}$ (topology on $\mathcal{V}$)	44

$\sim \mathcal{U} \sim$

uniformly continuous	76
Urysohn Lemma	58

$\sim \mathcal{V} \sim$

$\mathcal{V}$ (vertices)	3, 24, 39
compactness	47
variation	8, 17
design	8, 18
vector	
nonnegative	14
semipositive	15
strictly positive	14
vector field	
nonnegative	15
semipositive	15
strictly positive	15
vertex set	see $\mathcal{V}$
$\text{VF}(\mathcal{V})$ (variations)	17, 24, 60

$\sim \mathcal{W} \sim$

Whiteley, Walter	3, 14–36
wood	1
Wrinkle, Nancy	7

$\sim \mathcal{X} \sim$

$\mathcal{X}$ (design variations)	8, 18, 24
$\supset T(p)$	18
$\mathcal{X}$ -bar equivalent	8, 20
$\Leftarrow \mathcal{X}$ -infinitesimally rigid	20

$\Leftarrow$ strictly positive stress	9, 68
minimally	86
$\Rightarrow$ strictly positive stress	10, 86
partially	21
$\Leftrightarrow$ semipositive stress	5, 9, 65, 144
$\mathcal{X}$ -infinitesimally rigid	8, 20
$\Rightarrow \mathcal{X}$ -bar equivalent	20

$\sim \mathcal{Y} \sim$

Y (rigidity operator)	18, 24, 60
$Y(\mathcal{X})$ (loads)	19

$\sim \mathcal{Z} \sim$

$\mathbf{0}$ vs. 0	14
----------------------	----